

实证技术分析

用科学量化方法锁定交易信号

Evidence-Based Technical Analysis

Applying the Scientific Method and
Statistical Inference to Trading Signals

华章经典 · 金融投资

DAVID ARONSON

技术分析量化
实践经典

[美] 戴维·阿伦森 著

史雷 译

David
Aronson



机械工业出版社
China Machine Press

实证技术分析

Evidence-Based Technical Analysis

Applying the Scientific Method and
Statistical Inference to Trading Signals

华 章 经 典 · 金 融 投 资
D A V I D A R O N S O N

〔美〕戴维·阿伦森 著

史雷 译



机械工业出版社
China Machine Press

图书在版编目(CIP)数据

实证技术分析 / (美) 阿伦森 (Aronson, D.) 著; 史雷译. —北京: 机械工业出版社, 2015.2

(华章经典·金融投资)

书名原文: Evidence-Based Technical Analysis: Applying the Scientific Method and Statistical Inference to Trading Signals

ISBN 978-7-111-49259-7

I. 实… II. ①阿… ②史… III. 投资分析 IV. F830.593

中国版本图书馆CIP数据核字(2015)第023234号

本书版权登记号: 图字: 01-2013-0612

David Aronson. Evidence-Based Technical Analysis: Applying the Scientific Method and Statistical Inference to Trading Signals.

Copyright © 2007 by David R. Aronson.

This translation published under license. Simplified Chinese translation copyright © 2015 by China Machine Press.

No part of this book may be reproduced or transmitted in any form or by any means, electronic or mechanical, including photocopying, recording or any information storage and retrieval system, without permission, in writing, from the publisher.

All rights reserved.

本书中文简体字版由 John Wiley & Sons 公司授权机械工业出版社在全球独家出版发行。未经出版者书面许可, 不得以任何方式抄袭、复制或节录本书中的任何部分。

本书封底贴有 John Wiley & Sons 公司防伪标签, 无标签者不得销售。

实证技术分析

出版发行: 机械工业出版社(北京市西城区百万庄大街22号 邮政编码: 100037)

责任编辑: 黄姗姗

责任校对: 董纪丽

印 刷: 北京瑞德印刷有限公司

版 次: 2015年2月第1版第1次印刷

开 本: 170mm×242mm 1/16

印 张: 29.5

书 号: ISBN 978-7-111-49259-7

定 价: 75.00元

凡购本书, 如有缺页、倒页、脱页, 由本社发行部调换

客服热线: (010) 68995261 88361066

投稿热线: (010) 88379007

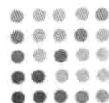
购书热线: (010) 68326294 88379649 68995259

读者信箱: hzjg@hzbook.com

版权所有·侵权必究

封底无防伪标均为盗版

本书法律顾问: 北京大成律师事务所 韩光/邹晓东



导 论

Introduction

技术分析的目的不仅是对金融市场数据中的重复模式进行研究,还要对市场未来的价格走势¹做出预测。技术分析包括很多种分析方法、形态、信号、指标以及交易策略,每一种理论的支持者都提出了各自的工作方法。

大部分传统或者普遍流行的技术分析在从以信仰为基础的民间艺术形式向以科学为基础的实践逐步发展的过程中,发挥着举足轻重的作用。也就是说,人们所讲的技术分析是通过形式各异的叙述和被大众所接受(最佳选择)的民间传说获得的,而不是依靠客观统计实证获得的。

本书的中心论点在于,技术分析必须要逐步发展成为一种缜密的观察性科学,才能够体现出这种理论的重要性。科学的方法是挖掘市场数据中有价值的知识,以及确定哪种技术分析方法具有预测力的唯一理性的方法。我把这种方法称为实证技术分析(EBTA)。实证技术分析以客观观察和统计推理(科学的方法)为基础,重新定义了被持续质疑的奇幻思维、容易上当受骗的痴迷者以及随机游走算法这三者之间的关系。

科学的方法并不适用于技术分析之外的其他分析法。科学的结论往往与那些直观明显的事情相冲突。例如,人们过去公认太阳围绕着地球运转,但是科学证明这种直觉是完全错误的。当金融市场行为两个最为显著的特征,即复杂性和高度随机性出现的时候,那些通过非正式、仅仅依靠直觉获取知识的方法就特别容易导致错误

VIII

的认知。尽管科学的方法并不能保证从市场数据这座大山中挖出金子，但是有一点可以确定，即非科学的方法只能挖到黄铜矿，而不是金矿。

本书的第二个中心论点是，包括流行的技术分析观点在内的大部分知识，并不能证明这种知识的合理性。

主要定义：论点与观点，信念和知识

我们总是使用**知识**和**信念**这两个术语，但是一直也没有对它们进行过严格意义上的定义。本书将会重复提到以上两个及其他的重要术语，所以在此要对它们进行正式的定义。

知识的基本组成部分是对事物的**陈述性表达**，也称为**观点和论点**。陈述句是4种表达方法之一，其余的还包括感叹句、疑问句和祈使句。陈述句之所以能够区别于其他的表达法是因为其本身包含真理的成分。换句话说，陈述句可以表示真实或者虚假的、大概真实或者大概虚假的意思。

例如：“超市中的橘子以5美分/打的促销价格出售”就是陈述句。这句话表达了在当地超市中某种事物现在的状态。它有可能是真实的，也有可能是虚假的。相反地，感叹句“天啊！太便宜了！”，祈使句“去，给我买一打”，或者疑问句“橘子是什么东西？”，都不能称为真实或者虚假的表述。

我们对技术分析的检验与陈述性表达有关，例如，“法则X具有预测力”。我们的目标就是确定哪些陈述性表达与我们的观点相符。

“我相信X”这句话是什么意思？“一般来说（‘事实’或者‘将会发生什么’），相信X意味着我们期望对X进行检验，前提是如果我们能够做到的话”。²因此，如果我们相信橘子以5美分/打的促销价格出售的观点，则意味着我去超市的话就有可能以5美分/打的价格买到橘子。然而，祈使句中的“去，给我买一打”，或者感叹句中对打折这个机会表现出的惊喜，都没有体现出以促销价格购买橘子的预期。

以上所提到的这些对于我们来说又意味着什么呢？我们可以把任何陈述都看作某一观点的后备选项，这些陈述必须“肯定某些预期事物的

状态”。³ 此类陈述被认为具有**可以认知的内容**，即它们表达某些人们**已经知道**的事情。“如果陈述当中不包含任何可以认知的信息，那么也就不存在能够让人相信的东西了”。⁴

并不是所有的陈述性表达都包含可以认知的内容。即使明显缺少认知性的内容，这也不会成为什么问题。例如：陈述句“星期二的平方根是素数”，⁵ 这种表达方式从字面上来讲，完全就是胡说八道（无稽之谈）。然而，其他的陈述性表达，缺少认知性的内容就不是很明显。这反而会出现一些问题，因为此类陈述会误导我们认为这种陈述中包含着预期的成分，实际上，陈述中并没有提及任何观点。这些虚假的陈述在本质上是**没有意义的观点或者空洞的论点**。

尽管没有意义的观点并不能够最终成为信仰，但这并不妨碍人们相信它们。例如，报纸上的星座专栏刊登的关于每天运程走势的模糊预测，以及宣扬虚假健康保健的人所做的承诺都属于没有意义的观点。由此可以看出，那些相信空洞的论点的人并没有意识到他们被告知的一切都是**没有认知性的内容**。

陈述中的认知性内容是否会成为一种信仰，判断的最好的方法就是美国物理学家霍尔（Hall）提出的**可识别的差异测试**⁶。“不论含有认知性内容的表达是真实还是虚假的，它们所表现出的差异都是显而易见的。这也就是包含信息的表达要比没有包含信息的表达更容易让人接受的原因”。⁷ 换句话说，根据可识别的差异测试的观点，如果陈述性的表达是真实的，我们就会对它抱有期待，如果陈述性的表达是虚假的，我们就会对事物的本质产生认知性的偏差。

存在明显差异的评判标准可以应用到对观点的预测当中。预测就是宣称知道未来将要发生的某些事情。如果某种预测包含认知性的内容，不论预测是否准确，最终的结果将会是非常明显的。许多时下流行的技术分析的从业者所做的预测都是缺少认知性的内容的。也就是说，使用技术分析的人所做出的预测太过模糊，以至于无法判断预测是否正确。

对于“橘子以 5 美分 / 打的促销价格出售”这个观点的真伪判断，只有当测试者亲自来到超市才能知道。正是因为产生了可以识别的差异，

我们便可以要求对这个观点进行检验测试。这一点我们会在第3章进行描述。对建立在可识别的差异基础之上的观点进行检验是科学的方法的核心。

霍尔在他的著作 *Practically Profound* 中，解释了他在对可识别的差异测试进行分析时，发现弗洛伊德的精神分析法是毫无意义的原因。

“弗洛伊德关于人类的性发展的观点与所有可能出现的情况是一致的。目前还没有任何一种方法可以对‘阴茎崇拜’或者‘完全阉割’的观点进行证明或者驳斥，因为我们无法对确认或者驳斥这种行为的解释做出明确的区分。完全相反的行为举止是同样具有预测力的，不论所谓的性心理是公开的还是受到抑制的”。“认知性内容的必要条件要将所有松散的、结构不合理的议论，或者过分坚持（阴谋论）在真实和虚假的事实之间不存在可以认知的差异的观点全部排除掉”。⁸ 在知识的脉络中，充满智慧的理论构想是不包括认知性的，从某种意义上来说，不论生命形态是怎样被发现的，它都会与由一些睿智的构想者⁹ 具体指出的潜在的生命形态的概念相一致。

什么是**知识**？知识可以定义为**经过证明是真实的信仰**。因此，为了使陈述性的描述符合**知识**的标准，不仅要求包含认知性的内容具有成为信仰的可能，而且还要满足其他两个条件。第一，它必须是真实的（或大概是真实的），第二，陈述必须具有让人信服的理由。信仰是在从实证进行推理的过程中，被证明是合理的存在。

早年人类持有一种错误的观点，即认为太阳围绕着地球转。很明显，这个观点是错误的。但是假设古人完全相信太阳围绕着地球运转是因为地球自转的原因，尽管这种观点是正确的，然而独立的个人也不能说就是有知识的。即使他们相信天文学家最终的结论是正确的，但是当时也没有证据去证实这个观点。在没有合理解释的情况下，就算是正确的观点也不能称为知识。以上这些概念请参见图 0-1。

伴随着前面所讲的知识，接下来我们来看一下**错误的信仰**和**虚假的知识**这两个概念。这两个概念都缺少知识需要具备的重要条件。因此，错误的信仰的出现不仅是因为它自身关注的是没有意义的观点，而且还

因为它尽管关注了一个有意义的观点，但是这个观点却没有经历过在实证中进行推理，并经过证明的过程。

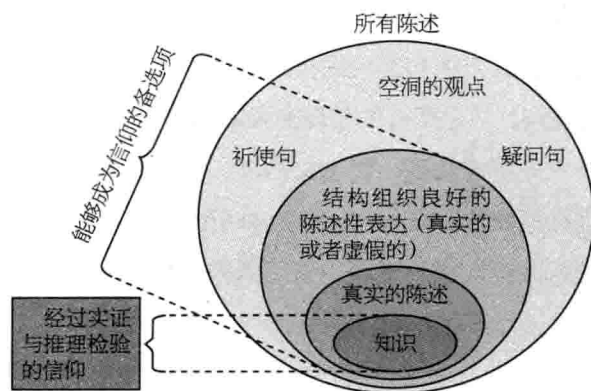


图 0-1 知识：经过证明是真实信仰

即使我们把每件事情都做到极致，并运用最有力的实证得出最好的结论，我们仍然可能接受错误的信仰。换句话说，如果我们从可以获得的大量证据当中得出合乎逻辑的有力结论，而这一结论也被证明是正确的，那么即使是谎言，我们也会相信，而且还会诚实地宣称我们知道这些事儿了。“在经过良好检验的实证网络中，当某一观点所描述的标的物能够经得起合理的怀疑，我们就有资格说‘我知道了’，但是即便是这样也不能保证我们知道全部”。¹⁰

当我们试图用观察实证的方法去了解世界的时候，谎言是我们在生活中无法逃避的事实。因此，基于科学的方法获得的知识与通过非正规方法获得的知识相比，虽然不确定的因素在减少，但是从本质上来讲依然存在着不确定性和临时性的特征。然而，随着时间的推移，科学的知识将会证明可以用一种循序渐进的、更加恰当的方法来描绘现实的世界。为了达到这个目的，我们任重而道远。实证技术分析的目标是最大限度地构成有关市场行为的知识体系，并且为实证和推理能力的积累设定范围。

错误的技术分析知识：不成熟的分析方法所付出的代价

时下普遍流行的技术分析方法所提出的观点往往无法让人信赖，要想知晓其中的缘由，我们必须把两种最典型的技术分析方法考虑进来，即：主观的技术分析方法和客观的技术分析方法。两种技术分析方法都会产生错误的观点，只是两者的表现形式不尽相同罢了。

客观的技术分析方法是一种定义清晰，并且能够对给出的明确的市场交易信号进行检验的分析方法。这种分析方法可以作为计算机算法和对历史数据进行回溯测试的工具。运用缜密的量化方法可以对通过回溯测试得出的结果进行评估。

主观的技术分析方法是一种没有经过明确定义的分析方法。因此，分析师就需要对这些模棱两可的概念做出解释。正是因为主观的技术分析方法的这种特性，使得它无法在计算机计算、回溯测试以及客观绩效评估中得到应用。也就是说，主观的技术分析方法并不能够对其自身的有效性进行判断，所以也就不会受到实证提出的质疑。

从实证技术分析的角度来看，主观的技术分析方法确实存在很多的问题。这些分析方法在本质上都属于没有意义的陈述，并且使人们对所表达的认知性内容产生错觉。这是因为主观的技术分析方法没有明确说明这些方法是如何应用的，所以分析师对同一组市场数据运用主观的技术分析方法就会得出不同的结果，进而也就无法对主观的技术分析方法能否提供有价值的预测进行判断。经典的图形形态分析、¹¹ 手绘趋势线、艾略特波浪原理、¹² 江恩形态、**Magic T's** 以及其他大量主观的技术分析方法都属于这一范畴。¹³ 主观的技术分析可以说是一种宗教信仰，再多的最优选择也不能表明这种方法可以成功解决其自身的问题。

主观的技术分析方法既缺少认知性的内容，也没有得到强有力的证据的支持，但是形态各异的主观的技术分析方法仍然不缺乏狂热的支持者。我们将在第 2 章解释在缺少证据或者面对自相矛盾的证据时，人们持有的坚定信念是如何通过错误的想法来产生的。

客观的技术分析方法同样会引起错误的观点，只是其产生的方式不

同而已。客观的技术分析方法对于从客观的证据中得出错误的推论是具有可追溯性的。事实上，在回溯测试中能够持续盈利的客观的技术分析方法并没有足够的证据证明它的优势所在。过往的业绩表现则会蒙蔽我们的眼睛。历史上的成功对于我们来说固然重要，但这并不代表我们就有充足的理由去推断一种分析方法具有预测力并且能够在未来盈利的可能。之前令人们交口称赞的业绩很可能靠的是运气的成分，或者是由于一种被称为数据挖掘的回溯测试所产生的上涨型偏差。对回溯测试中产生的盈利是根据好的分析方法，还是依靠运气的判断，唯一的方法就是通过缜密的统计学推断。这些内容我们将在第4章和第5章进行讨论。第6章关注的是数据挖掘偏差的问题。我坚信准确运行的数据挖掘是当今技术人员发现新知的最佳方法，而特殊的统计测试也将用于获取挖掘结果的工作中。

实证技术分析的与众不同之处

究竟是哪些因素使得实证技术分析有别于其他普遍流行的技术分析呢？首先，实证技术分析对含有意义的陈述做出了预测——因为客观的技术分析方法能够通过历史数据进行检验。其次，实证技术分析利用先进的统计推理方法来确定回溯测试的可盈利性是否具备有效的方法的特征。因此，实证技术分析最为核心的焦点在于决定哪种客观的技术分析方法具有实操性。

实证技术分析排斥任何形式的主观的技术分析方法。主观的技术分析方法比错误的技术分析方法还要糟糕。之所以把某种陈述确定为错误的（不真实的），是因为陈述中所传递的认知性内容是没有经过检验的，即主观的技术分析方法没有提供任何认知性内容。从表面上看，主观的技术分析方法确实也提供了一些认知性的概念，但是一旦它们经过严格的检验，其空洞的陈述的本质就显露无遗了。

新世纪健康疗法的倡导者就非常善于空洞的陈述。他们会告诉你戴上一种神奇的铜手镯可以使你神清气爽，走起路来步履轻盈，还可以提高你的高尔夫球的成绩，甚至对你的爱情也有帮助。然而，这种缺少具

XIV

体内容的陈述是不可能实现它的承诺的，更不用说去检验它的真伪了。空洞的陈述是永远无法用客观的证据（事实）进行证明或者反驳的。根据以上的推理，可以说主观的技术分析的观点是空洞的，是无法接受实证检验的。因为技术分析必须是值得人信赖的。

相反地，有含义的陈述是可以检验的，因为它提供了可以进行检验的内容。它说明了怎样才能提高你的高尔夫球的成绩，以及为什么你在走路的时候会步履轻盈。这种明确的陈述可以通过实证证据进行反驳。

从实证技术分析的角度来看，主观的技术分析方法的拥护者面临着这样一个选择：要么重新制定客观的技术分析方法，像艾略特波浪原理的实践者所做的那样，¹⁴ 让这种方法易于接受实证的检验和反驳；要么他们必须证明这种新方法是值得依赖的，或许江恩线实际上会提供有价值的信息，但是以它目前的形态，我们还是拒绝承认这种知识。

至于客观的技术分析方法，实证技术分析没有对可盈利性进行表面上的回溯测试。取而代之的是根据缜密的统计学评估来判断盈利的出现是源自运气，还是源自有偏见的研究。我们在第 6 章还会提及这一点，在众多案例当中，进行可盈利的回溯测试只能使数据挖掘者挖到黄铜矿（假的金矿）。这就可以解释为什么在回溯测试中表现上佳的客观的技术分析方法要比应用到新数据中的客观的技术分析方法的表现要差了。实证技术分析运用的密集型计算机的统计学方法将数据挖掘偏差中出现的问题最小化了。

从技术分析到实证技术分析的发展也包含了道德因素。所有分析师都要遵守道德和法律的责任与义务。无论你运用哪种分析方法，最后提出的建议都要基于合理的基础，而不能成为无理可依的空洞的陈述。¹⁵ 确认一种分析方法是否有价值的唯一合理的基础就是客观的实证。主观的技术分析方法是达不到这一标准的。客观的技术分析的运用与实证技术分析的应用标准是一致的。

从学术的角度产生的实证技术分析

实证技术分析不是一个全新的概念。在过去的 20 年中，权威的专

业期刊中¹⁶就出现了本书所提倡的使用一种缜密的技术分析方法的观点。¹⁷有些研究显示技术分析是不起作用的，而有的研究又说它是有作用的。因为每一种研究都局限于技术分析某一个特定的方面，或者是具体的数据群，这就使得每种研究都会得出不同的结论。这种现象在科学当中是相当普遍的。

以下这部分是从专业的技术分析中得到的研究结果。它显示了用缜密理智的方法，证明技术分析是值得研究的领域。

- 专业的图表分析专家不能够从随机过程¹⁸产生的趋势图中分辨出实际的股票市场价格走势图。
- 关于商品¹⁹和外汇市场的实证证据可以通过利用简单的客观趋势指标获得。此外，趋势跟风的投机者所获得的利润，可以通过经济理论²⁰进行解释。因为他们的行为给商业套期保值者提供了有价值的经济服务，因此，价格的风险也从套期保值者一方转移到了投机者一方。
- 当简单的技术法则应用于对由相对年轻的公司组成的股票市场（罗素2000指数，纳斯达克综合指数）²¹的平均价格的研究时，单独或者综合应用简单的技术分析法则可以在统计以及盈利上带来巨大收获。
- 神经网络可以将简单移动平均法则中显示的买/卖信号与非线性模型结合起来。二者的结合显示了对1897～1988年的道琼斯平均指数进行预测时产生的良好效果。²²
- 通过简单的动量指标检测到行业集团和产业集群的趋势之后，选择长期持有的策略，进而赚取额外的利润。²³
- 股票市场显示了之前市场相对强弱的程度，并且继续显示在未来3～12个月的水平期间内高于和低于平均线的业绩表现。²⁴
- 我们在接近52周高点的位置卖出美国股市中的股票，此时的业绩表现远远超过其他股票市场的股票。以股票的现价与其52周高点之间的差额作为指标，对未来的相对绩效起到了非常有价值的预

测作用。²⁵ 这个指标对于预测澳大利亚股市的股价更加有效。²⁶

- 当用一种客观的方式对货币市场进行检验时，头肩形态就限制了预测力的发挥，此时用简单的过滤法则可以得到更好的结果。当头肩形态用于对股票市场的客观检验时，就不能提供有用的信息了。²⁷ 交易者根据这个信号进行操作，等同于跟踪随机信号进行的操作。
- 通过股票的交易量对有价值的预测信息²⁸ 进行统计，可以提高以跟踪公告所产生的价格的大幅变化为基础发出的信号的收益率。²⁹
- 密集型计算机数据建模神经网络、遗传算法以及其他的统计学知识和人工智能方法已经根据技术指标³⁰ 发现了盈利模式。

我是在扮演批判技术分析的角色吗

1960 年，也就是我 15 岁的时候，我对技术分析产生了兴趣。在我上高中和大学期间，我运用绘制点数图的方法跟踪了大量的股票。我从 1973 年开始使用专业的技术分析方法，起初，我作为股票经纪人，随后又作为一家小型软件公司——兰登研究集团有限公司的主管合伙人，那个时候我是最先接受在金融市场中使用机器学习和数据挖掘的人。后来我又为 Spear, Leeds & Kellogg³¹（简称 SLK 公司）担任业主权益交易商。1988 年，我获得了注册市场技术分析师的资格。我个人在技术分析方面的藏书超过 300 本，发表论文十几篇，并且就技术分析这一课题进行过多次演讲。目前，我在纽约市立大学巴鲁克学院席克林商学院讲授技术分析的研究课程。我必须承认我以前的文章和研究并没有达到实证技术分析的标准，特别是在统计的显著性和数据挖掘偏差方面。

鉴于 5 年中我为 SLK 公司进行商业资本运作所取得的平淡业绩，我长期以来对技术分析所抱有的坚定信念发生了动摇。为什么我如此信任的技术分析会表现得如此差劲呢？到底是我个人的原因还是其他原因呢？我在哲学方面受过的专业训练为我对技术分析不断出现的质疑提供了有力的证据。在我读完托马斯·吉洛维奇（Thomas Gilovich）的《理性犯的错》（*How We Know What Isn't So*）以及迈克尔·谢尔梅尔（Michal

Shermer) 的《为什么人们总是相信荒诞怪异的事情》(*Why People Believe Weird Things*) 两本书之后, 我对技术分析产生了完全的质疑。我自己的结论是: 包括我本人在内, 对技术分析的本质理解并非我所想的那样, 我也在相信一些荒诞、怪异的事情。

技术分析: 艺术、科学还是迷信

在技术分析界有这样一种争论: 技术分析是艺术还是科学? 这个问题本身的提法就有问题, 它应该这样表述才对: 技术分析是以科学还是迷信为基础的? 如果是这种问法, 争论也就不存在了。

有人会说, 以科学的、可以检验的方法对技术分析这一知识概念进行解释的时候, 这种方法本身就包含着太多的差异。我对此的回应是, “技术分析是不可检验的” 这句话听上去很像一种知识, 但实际却不是这样。技术分析是属于天文学、数字占卜术以及其他非科学实践的迷信行为。

创造性和灵感在科学中发挥着重要的作用。二者同样适用于实证技术分析。所有对科学的探寻都是从假设开始的。一个新的想法或者观点都是以从前的知识、经历和直觉组成的神秘结合体作为灵感来源的。然而, 好的科学和严格的分析方法保持创造性。自由地提出新的想法必须与严格的行为准则相结合, 进而淘汰那些在客观测试的严格考验中被证明是无效的想法。如果不这样做的话, 当人们异想天开的时候, 奇幻的思维就会取代严谨的思想。

技术分析试图运用精确的物理定律进行预测, 并以此发现新的法则是不可能的。金融市场中固有的复杂性和随机性, 以及受约束的实验阻碍了研究结果的产生。然而, 精确的预测并不是科学定义的必要条件。相反地, 它是通过明确的定义来认知并淘汰那些错误的想法的。

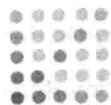
对本书我有 4 个愿望: 第一, 鼓励技术分析师之间的对话, 并最终把我们的专业领域建立在牢固的知识的基础之上; 第二, 鼓励根据此处提供的线索继续探索; 第三, 希望技术分析的使用者能够从以技术分析为基础提供产品和服务的人那里汲取更多的知识; 第四, 希望专业或者

XVIII

非专业的技术分析实践者明白他们在人机交互中的重要作用。通过人机交互，我们可以加速合理的技术分析知识的发展。

毫无疑问，技术分析的从业者并不认同这些说法。这的确是件好事。因为进入牡蛎身体里的细沙有时也会变为珍珠。我希望我的同事继续发挥他们的能量，去寻求更加合理的知识，而不是为那些站不住脚的谬论做无用的辩护。

本书由两部分组成。第一部分建立了实证技术分析在方法论、哲学、心理学以及统计学方面的基础。第二部分演示了实证技术分析的一种方法，即引用 25 年来的历史数据对用于标准普尔 500 指数的 6 402 个二进制买 / 卖法则进行测试，并且对用于测试旨在处理数据挖掘偏差问题的法则的统计显著性进行评估。



致 谢

Acknowledgements

虽然我们应当把一本书的成功归结于它的作者，但实际上这却是很多人共同努力的结果。在此我要对他们表示感谢，没有他们的付出和劳动，这本书不会如此顺利地出版。

我要感谢 **Timothy Masters** 博士，在我们相识的这十多年里，他的耐心和对我的引导让我坚定地走上了统计学分析这条道路。**Tim** 不仅在技术问题方面对我产生重要的影响，而且还负责编写代码及运行 **ATR** 法则试验，并且对进行测试的 6 400 多个法则进行统计学检查。**Tim** 还发明了蒙特卡罗置换法，用于替代怀特的真实检验 (**White's Reality Check**) 方法，他用该方法对通过数据挖掘发现的法则的统计显著性进行测试。同时，**Tim** 还决定把这种新方法推广到大众领域，并且通过本书在第一时间和公众见面。

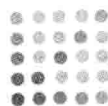
我还要感谢 **Stuart Okorofsky** 和 **John Wolberg**，前者在编程方面显示出了过人的天赋，后者是数据库创建方面的专家。此外，我还要感谢怀特的真实检验方法的创始人——**Halbert White** 博士以及马萨诸塞大学阿姆赫斯特分校知识探索实验室的主任戴维·詹森 (**David Jensen**) 教授所提供的帮助。

最后，我还要对为本书各个章节进行检查和注释的工作人员表示最诚挚的感谢，他们带给我的反馈意见非常重要：他们是 **Charles Neumann**, **Lance Rembar**, **Dr. Samuel Aronson**, **Dennis Katz**, **Hayes Martin**, **George Butler**, **Dr. John Wolberg**, **Jay Bono**, **Dr. Andre**

XX

Shlefier, Dr. John Nofsinger, Doyle Delaney, Ken Byerly, James Kunstler,
and Kenny Rome.

特别感谢 John Wiley& Sons Inc 公司：感谢 Kevin Commins 发现
了对技术分析进行严格评估的重要性，以及 Emilie Herman 为本书的编
辑工作所付出的努力。在这里对 Michael Lisk 和 Laura Walsh 一并表示
感谢。



目 录

Contents

导 论

致 谢

第一部分

方法论、心理学、哲学以及统计学的基础

第1章 客观的法则及其评估 / 2

重要的分水岭：客观的技术分析 VS. 主观的技术分析 / 2

技术分析法则 / 3

传统的法则与反转法则 / 8

在法则评估中基准的使用 / 9

其他细节：前视偏差和交易成本 / 15

第2章 主观的技术分析的效率错觉 / 18

主观的技术分析作为知识的不合理性 / 20

个人的传闻轶事：从最初真正的技术分析信仰者到后来的怀疑者 / 21

大脑：自然形态的发现者 / 24

荒诞的信仰的传播 / 26

认知心理学：启发、偏好和错觉 / 27

人类处理信息的局限性 / 28

极端的确认：过度自信偏差 / 30

IV

	二手信息偏差：好故事的力量 / 44
	确认性偏差：现存的信念是如何过滤经验的，以及矛盾的实证是 如何存活下来的 / 48
	虚幻的相关性 / 58
	图形分析中的错误信仰 / 68
	直觉判断与启发的作用 / 72
	代表性启发式与错误趋势和图形中的形态：真实与虚幻 / 79
	解决虚幻的知识的方法：科学的方法 / 87
第3章	科学的方法与技术分析 / 88
	最重要的知识：获得新知的方法 / 88
	希腊科学的遗产：喜忧参半的结果 / 89
	科学革命的起源 / 90
	对客观现实的信心与客观观察 / 92
	科学的知识的本质 / 93
	逻辑在科学中的作用 / 96
	科学的哲学 / 110
	最终结果：假设演绎法 / 132
	对观察结果进行严谨和详细的分析 / 135
	对科学的方法的主要内容的总结 / 136
	如果技术分析采用科学的方法 / 137
	主观的技术分析客观化：样本 / 140
	技术分析的子集 / 150
第4章	统计分析 / 152
	统计学推理概览 / 153
	严谨的统计学分析的必要性 / 159
	抽样与统计推断的样本 / 160
	实验概率与随机变量 / 162
	统计理论 / 174
	描述性的统计学 / 178
	概率 / 181

	随机变量的分布概率	/ 184
	概率和概率分布的部分面积之间的关系	/ 186
	抽样分布：统计学推理当中最重要的概念	/ 188
	获得抽样分布：经典的方法	/ 194
	用计算机模拟计算的方法获得抽样分布	/ 200
	关于下一章	/ 201
第5章	假设检验和置信区间	/ 202
	统计学推理的两种类型	/ 202
	假设检验 VS. 非正式推理	/ 203
	假设检验的基本原理	/ 207
	假设检验：构成法	/ 212
	用计算机模拟的方法产生抽样分布	/ 218
	估算	/ 227
第6章	数据挖掘偏差：客观技术分析的黄铜矿	/ 236
	落入陷阱：数据挖掘偏差的故事	/ 237
	在客观的技术分析中出现的错误知识的问题	/ 243
	数据挖掘	/ 246
	客观的技术分析研究	/ 250
	数据挖掘和统计推断	/ 254
	数据挖掘偏差：两种因素的影响	/ 260
	数据挖掘偏差的试验研究	/ 272
	解决方法：处理数据挖掘偏差	/ 299
第7章	非随机价格波动理论	/ 311
	理论的重要性	/ 311
	科学的理论	/ 312
	流行的技术分析有什么问题	/ 312
	反对者的观点：有效市场和随机游走	/ 314
	挑战有效市场假说	/ 322
	行为金融学：随机价格波动理论	/ 337

发生在有效市场条件下的非随机价格运动 / 360

结语 / 368

第二部分

案例研究：标准普尔 500 指数的信号法则

第8章 对应用于标准普尔500指数的法则进行数据挖掘的案例

研究 / 370

数据挖掘偏差和法则评估 / 370

避免数据探测法偏差 / 371

分析数据序列 / 372

技术分析的主题 / 372

绩效统计量：平均收益率 / 372

不评估复杂法则 / 373

以统计学术语定义的案例研究 / 373

法则：将数据序列转换成市场头寸 / 375

时间序列指标 / 377

法则的输入序列：原始时间序列和指标 / 384

应用于案例研究的 40 种输入序列列表 / 395

法则 / 396

第9章 案例研究结果与技术分析的未来 / 415

研究结果 / 415

对案例研究的批评 / 422

案例研究扩展的可能性 / 426

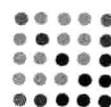
把复杂法则考虑在内 / 427

技术分析的未来 / 435

附录 对消除趋势与以头寸偏差为基础的基准相等的观点进行 论证 / 448

注释^①

① 本书注释请见华章公司网站：<http://www.hzbook.com>。



第一部分 Part 1

方法论、心理学、哲学以及 统计学的基础



第1章

Chapter 1

客观的法则及其评估

本章介绍的是客观的二元信号法则的概念，以及对其进行严谨评估的方法论。它在不包含充足信息的信号的盈利性的基础上，对评估基准进行了定义。本法则也体现了通过消除市场数据的长期趋势，与持有不同多空头寸偏好的法则所产生的绩效进行对比的必要性。

✓ 重要的分水岭：客观的技术分析 VS. 主观的技术分析

从广义上来讲，技术分析（technical analysis, TA）分为两大类，即客观的技术分析和主观的技术分析。主观的技术分析包括不能够明确定义的分析方法和图形形态。其结果就是从主观的方法当中得出的结论反映了分析师对应用这种方法所做出的私人解释。这就有可能使得两位分析师运用同一种方法对一组市场数据进行分析，而得出完全不同的结论。由于主观的分析方法是不可检验的，也就使得其自身的有效性无法得到实证的检验。因此，这就为这种错误的观点的繁荣发展提供了肥沃的土壤。

反过来，客观的方法是可以明确定义的。当把一种客观的分析方法应用到市场数据中时，它所提供的信号和预测是非常清晰的，这就使得对历史数据以及判断绩效的准确程度方面的模拟成为可能。我们把这种方法称为回溯测试。对一种客观的方法进行的回溯测试是可以重复实验的。这种实验可以通过统计数据对盈利状况进行测试，并且对客观的方法进行反驳。回溯测试可以发现哪种客观的方法是有效的，哪种是无效的。

酸性试验是从主观的方法当中分辨出客观的方法的程序化准则：假设该方法可以被构建为能够产生明确的市场头寸（长期、¹短期²或中性³）

的计算机程序，那么就属于客观的分析方法。所有方法不能够被主观地、默认地归纳于这种程序之中。

✓ 技术分析法则

客观的技术分析也被称为机械的交易法则或者交易系统。在本书中，所有客观的技术分析方法全都简称为法则。

法则就是指将信息中的一项或者多项（法则的输入量），转变为符合人们预期的市场头寸（长期、短期、中性，即法则的输出量）的函数。输入量包含一组或者多组金融市场的时间序列。而法则是由一个或者多个数学、逻辑运算符号定义的。它通过把输入的时间序列转换成为一组包含连续市场头寸（长期、短期、场外）的新时间序列。输出则是用带有正负号的数字（例如 +1, -1）来表示的。本书采用正值表示多头头寸（long position），用负值表示空头头寸（short position）。法则是把一项或者多项信息的输入序列转换成为输出序列的过程（见图 1-1）。

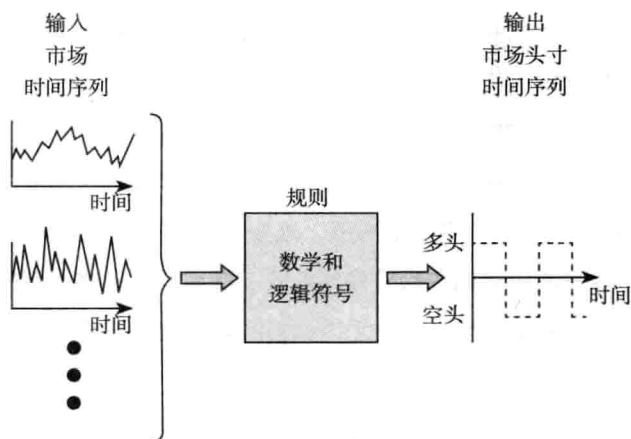


图 1-1 技术分析法则将输入时间序列转换成为市场头寸时间序列

当输出时间序列的数值发生改变的时候，法则就会发出信号，或者说当之前已经设定好的市场头寸发生改变时，信号就会出现。例如，输出值从 +1 转变为 -1 就说明了之前持有多头头寸的结束，以及新的空头头寸的开始。输出值也可以不设定为 (+1, -1)。在复杂的法则当中，也可以设

定规模大小不同的市场头寸，输出值的范围可以是 $(+10, -10)$ 。例如，输出值为 $+10$ 代表有10手多头头寸得到确认，例如，10手多头铜期货合约。输出值从 $+10$ 变为 $+5$ 则表示多头合约从10手减少到了5手（卖出5手）。

二进制法则和临界值

包含二进制输出的法则是最简单的法则。换句话说，它的输出值仅可以假设为两个数值，即 $+1$ 和 -1 。二进制法则也可以用于设定多头/中性头寸，或者空头/中性头寸。本书中所指的法则是指二进制多头/空头 $(+1, -1)$ 。

以二进制多头/空头头寸为基础制定的投资策略通常是指市场交易当中的多头或空头。这种法则指的是**反转法则**，因为交易信号标志着多头向空头，或者空头向多头的反转。随着时间的推移，反转法则产生了 $+1$ 或 -1 的时间序列。这代表了多头和空头的交错序列。

用于定义法则的明确的数学和逻辑符号可以产生非常大的变化，然而它们之间也存在着共同点。首先，是**临界值**的概念，它是指把输入时间序列中信息变化的临界水平从不相关的波动中分辨出来。前提是输入时间序列必须是信息与噪声的结合体。

当时间序列突破临界值的时候，使用临界值的法则就会发出信号。这个数值可能高于，也可能低于临界值。这些关键点可以用**大于号** $(>)$ 或者**小于号** $(<)$ 这种不等式的逻辑符号来检验。例如，如果时间序列大于临界值，则该法则的输出值为 $+1$ ，反之，如果时间序列小于临界值，则该法则的输出值为 -1 。

临界值可以设定为一个固定的值，也可以随着用于分析的时间序列变化的结果而变化。可变的临界值适用于显示趋势的时间序列，而这一趋势在时间序列的层面上是持续变化的。固定的临界值法则在此是不适用的，通常情况下将其应用于资产定价（例如，标准普尔500指数）和资产收益（AAA级债券收益）。移动平均线与亚历山大反转过滤器（Alexander reversal filter）都属于锯齿型过滤器（zigzag filter），它们都是时间序列符号

的样本，通常用于对可变临界值的定义。在本书的法则中运用的运算符号将在第8章中进行详解。

关于可变的临界值是如何在显示趋势的时间序列中发出信号的，移动平均交叉法则可以说是最好的实例。当时间序列线与移动平均线相交时，移动平均交叉法则就会发出信号，例如：

如果时间序列线在移动平均线之上，则法则的输出值为+1。反之，如果时间序列线在移动平均线之下，则法则的输出值为-1。

请参见图1-2。

由于使用了单一的临界值，通过移动-平均-交叉法则所发出的信号，从定义上来说，是相互排斥的。因为只提供了一个单一的临界值，其结果就只有两种可能，即时间序列线在临界值之上或者之

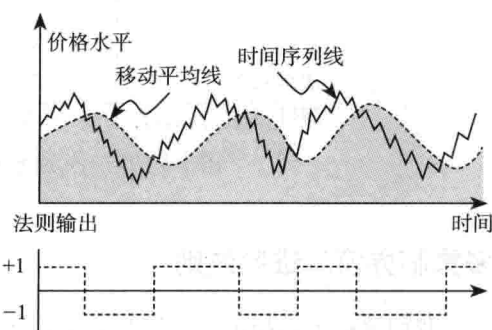


图1-2 移动-平均-交叉法则

下。⁴ 这种情况是显而易见的（不存在其他的可能性）。⁵ 因此，该法则所发出的信号是不可能不一致的。

使用固定的临界值适用于没有显示趋势的市场时间序列。我们把这种时间序列称为**平稳的时间序列**。对于平稳的时间序列，是有着严格的数学定义的，但是我在这里所运用的术语是一个较为松散的概念，意思是说随着时间的推移，一条时间序列有一个相对稳定的数值，而它的波动范围也被限制在一个大致的水平范围之内。技术分析的实践者通常把这种现象称为**振荡器**。

时间序列不仅可以显示趋势，也可以**消除趋势**。换言之，即时间序列可以转换为一种平稳的主序列。消除趋势这一概念将在第8章中详细介绍，它通常涉及差异或者比率的情况。例如，时间序列与移动平均线的比率会产生出一条原始时间序列的变体。一旦消除趋势形成，时间序列就会

围绕着一个相对稳定的均值在一个界线相对清晰的水平范围内进行波动。一旦时间序列已经形成平稳的态势，固定的临界值法则就可以发挥作用了。图 1-3 中所显示的是固定的临界值法则运用一个数值为 75 的临界值的例子。当时间序列大于临界值时，法则的输出值为 +1，当时间序列小于临界值时，法则的输出值为 -1。

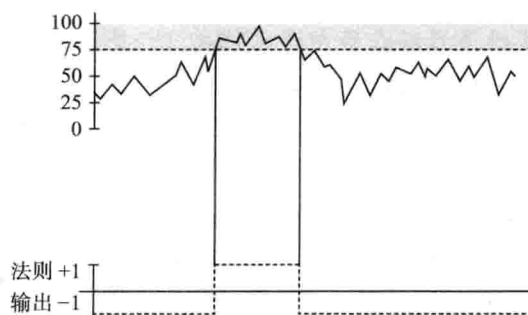


图 1-3 单一固定的临界值法则

多重临界值二进制法则

我们之前曾经指出，二进制法则是从单一临界值中衍生出来的，因为临界值对两种互相排斥的情况给出了定义，即时间序列要么在临界值之上，要么在临界值之下。然而，二进制法则依然可以通过多重临界值的方式衍生出来，但是应用一种以上的临界值会形成输入时间序列出现超过两种情况的可能性。其结果是，多重临界值需要比简单的不等式运算符（大于或者小于）更为复杂的逻辑运算符，而简单的不等式运算符是完全可以满足单一临界值的需求的。

当出现两个或者更多的临界值时，会出现不止两种的可能情况。例如，当出现上端和下端两个临界值时，对输入时间序列来说会出现 3 种可能性，它可能会出现在临界值之上，也可能出现在临界值之下，还可能出现在两个临界值之间。要想在这种情况下建立二进制法则，可以用两个互相排斥的事件来定义该法则。通过时间序列在固定的方向中与固定的临界值交叉来定义第一个事件。因此，一个事件触发了某种法则的输出值，而这个输出值会一直持续到第二个事件的发生，而这第二个事件又是与第一个事件

完全背离的，然后再接着触发另外一个输出值。例如，上端临界值的向上交叉会触发 $a+1$ ，而下端的临界值向下的交叉会触发 $a-1$ 。

运用这种法则的逻辑运算符可以用触发器这个词来替代。触发器这个词来源于当出现某一事件的时候，法则的输出值向一个方向移动的事实。然后，当第二个事件发生的时候，输出值会向着相反的方向移动。“触发器”的逻辑可以运用于可变临界值或者固定临界值法则。基于两个可变临界值的法则，可以参照移动平均线法则（见图 1-4）。移动平均数是围绕着上端线或者下端线运动的。这条线可以是在移动平均数之上或者之下的固定百分比，也可以像布林线（Bollinger Band）⁶ 一样，在时间序列最近波动的基础之上，

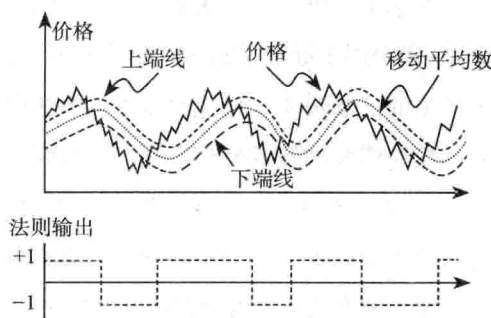


图 1-4 移动平均线法则

产生出与布林线的变化相背离的状态。数值为 +1 的输出变量由于上端临界值的向上突破而被触发。这个数值会一直持续下去，直到下端的临界值被突破的时候为止，此时，数值会从 +1 向 -1 转变。

显然，我们还可以想到很多其他的可能性。我们现在所展现出来的，是可以将时间序列转变为可以显示市场头寸的时间序列。

海斯（Hayes）⁷ 利用定向模式对临界值规则做了补充。他把多重临界值应用到诸如扩散指标⁸ 这类比较稳定的时间序列当中。在给定的时间点上，指标的模式由它所占据的区域和最近走势的变化（最近 5 周的上升或下降走势）确定。每一个区域由上端或者下端的临界值（40 和 60）确定。海斯将这种方法应用到一种名为“Big mo”的专用扩散指标中。通过两个临界值和可能的定向模式（上升或下降），6 种相互排斥的情况就可以被确定了。通过把一种输出变量（-1）分配到这 6 种情况当中的分析，一个二进制规则就可以从中衍生出来，然后，再把另一个输出变量（-1）分配到另外的 5 种可能性当中。海斯声称，当扩散指标在 60 以上时，就是上涨的趋势，而股票市场的收益率（价值线综合指数）是 50%。这种情况在

1966 ~ 2000 年的出现概率是 20%。然而, 当扩散指标大于 60, 且其最近的走势是消极的, 那么市场的年收益就是 0。这种情况在 1966 ~ 2000 年的出现概率是 16%。⁹

传统的法则与反转法则

本书第二部分的主要内容是案例研究, 即引用 25 年的历史资料测试 6 400 多个用于标准普尔 500 指数的多 / 空二进制法则的绩效。其中很多法则产生的市场头寸与传统的技术分析法则生成的市场头寸是相一致的。例如, 按照传统的技术分析法则, 当时间序列在移动平均数之上的时候, 移动平均交叉法则将其解读为牛市 (输出值为 +1), 而当时间序列在移动平均数之下的时候, 移动平均交叉法则将其解读为熊市 (输出值为 -1)。我个人把这种方法称为**传统的技术分析法则**。

传统的技术分析的真实性是值得怀疑的, 我们需要对与传统的解释相悖的法则进行检验。换句话说, 按照传统的假设, 对某一形态做出价格上涨的预测, 结果很有可能是该种形态呈下跌的走势。或者说, 没有任何一种形态具有预测的价值。

我们可以通过创建一组额外的法则来实现这一问题, 只要这种法则的输出值与传统的技术分析法则相反就可以了。我所说的是**反转法则** (见图 1-5)。当输入时间序列在移动平均数之上, 移动平均交叉法则的输出值会反转成为 -1。当输入时间序列在移动平均数之下, 其输出值为 +1。

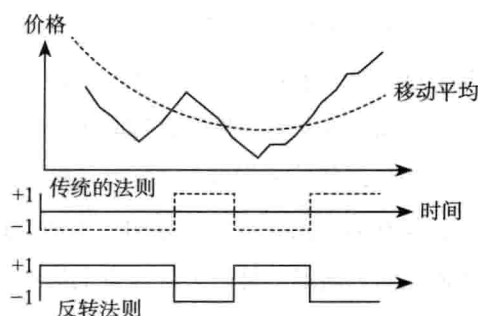


图 1-5 传统的法则与反转法则

我们还可以从其他的方面来考虑反转法则。许多在第二部分用于测试的法则, 使用的是输入序列, 而不是标普 500 指数, 例如, BAA 级和 AAA 级公司债券的收益率是不同的。这种时间序列是如何发出信号的, 到

目前为止还不是很明显。因此，上涨趋势与下跌趋势在收益率方面存在的差异可以被认为是买入信号的不同。具体的细节我们将在第8章中进行讨论。

在法则评估中基准的使用

对于很多领域来说，绩效是相对的。也就是说，是相对于基准的表现。这个基准不是纯粹的绩效标准，而是非正式的。在田径场上，投掷铅球的运动员的成绩是以当天的最好成绩，或者世界纪录作为基准的。比如说某个运动员的成绩是43英尺，^①我们并不能因此来判断他的成绩如何，但是当我们知道之前最好的成绩是23英尺的时候，43英尺就是非常了不起的成绩了！

基准是与法则的评估有关系的。当绩效数字与相关的基准相比较时，它们会提供有价值的信息。例如，一种法则的收益率是10%这一事实，其在回溯测试中是毫无意义的。如果说在这组数据中，其他的法则能够获得超过30%的收益率，那么收益率为10%的这个数据就显得太差劲了，相反地，如果其他的法则所获得的收益率微乎其微，那么这10%的收益率简直堪称完美。

对于技术分析的绩效来说，有没有恰当的基准呢？要想证明一种法则是否有效，我们要以什么样的标准来衡量呢？其实适当的标准有很多。本书把不具有预测力（随机产生的信号）的法则的绩效定为衡量的标准。这个方法与在其他领域应用的科学实践有异曲同工之处。在医学界，一种新药必须具有让人信服的功效，比如安慰剂（糖丸）是能够证明其是有用的。当然，理性的投资者也许会理智地选择更高的绩效标准，而不是低的标准。其他易于理解的基准是：无风险收益率、买入并持有策略的收益率以及目前正在使用的法则的收益率。

实际上，没有哪种法则比基准更有说服力。要想证明某种法则是优越的，必须先把盈利是靠偶然（好运气）实现的这种可能性排除出去。在一

① 1英尺=0.3048米。

个由运气成分组成的数据样本中，不具有预测力的法则是完全可以击败衡量它的基准的。将依靠运气获得的盈利排除在外，可以对统计的显著性给出解释。这一点我们将在第4～6章讨论。

我们已经创建了将要使用的基准，而这个基准带来的收益率是通过不具有预测力的法则获得的。现在我们要面对另外一个问题：没有预测力的法则能够产生多少收益率？表面上看，对其合理的预期是0。然而，这种结果只会出现在特殊的、受到限制的条件下。

实际上，对没有预测力的法则的收益率预期结果显然不是0。这是因为一个法则的绩效会受到与其预测力无关的因素的影响。

头寸偏好的协同效应与市场趋势在回溯测试中的表现

实际上，一种法则在回溯测试中的绩效是由两个独立的部分组成的。其一，是法则的预测力（如果有的话）。这是我们感兴趣的内容。其二，绩效的组成部分是通过与法则的预测力无关的两个因素得出的：①法则的多/空头寸偏好，②在进行回溯测试期间市场的净趋势。

绩效的因素可以影响回溯测试的结果，并且造成对法则进行预测的困难。它可以使一个没有预测力的法则产生数值为正的 average 收益率，或者使一个本身具有预测力的法则产生数值为负的平均收益率。除非将绩效的内容移除，否则是不可能得到对法则的精确评估的。我们现在来分析这两种因素的内容。

第一个因素是法则的多/空头寸偏好。即在回溯测试中，法则在 $a+1$ 输出状态下的时间与在 $a-1$ 输出状态下的时间的比例。如果其中一个输出状态在回溯测试中处于主导地位，那么我们就可以说这个法则具有多头偏好。例如，如果输出状态长期处于多头状态，则该法则就具有多头偏好。

第二个因素是市场净趋势，或者是在回溯测试中市场的平均每日收盘价的变化。如果市场的净趋势不是零，这就说明法则具有多/空头寸偏好，而该法则的绩效也会因此受到影响。换句话说，把不符合要求的绩效从通过法则的实际预测力产生的绩效中增加或者减少，会歪曲回溯测试的结果。然而，如果市场的净趋势等于零，或者说该法则没有头寸偏好，则法则过

去所获得的利润就完全依靠法则的预测力（增加或减少随机变量）。这一点我们将在后期通过数学的方式进行证明。

试想一下，一种技术分析法则有多头偏好，但是我们又知道其并不具有预测力。对于这个法则所发出的信号，我们通过轮盘赌来进行模拟。为了创建多头偏好，轮盘中的大多数卡槽会被设定为多头（+1）。假设 100 个卡槽以如下的方式进行划分：75 个为 +1，25 个为 -1。每一天，通过转动轮盘来决定当天是多头还是空头。如果在这期间市场的平均每日收盘价（average daily）大于零（例如，净趋势是上涨趋势），那么即使这个法则不包括有预测性的信息，那么也会产生数值为正的预期收益率。一种法则的预期收益率可以通过用于计算随机变量的预期值的公式来计算（后面详细讨论）。

就像一个不具有预测力的法则可以产生正值的收益率那样，一个具有预测力的法则同样可以产生负值的收益率。当一种法则的头寸偏好与市场趋势相背离时，这种现象就会出现。把市场趋势的影响与法则的头寸偏好相结合，就可以将由法则的预测力所导致的正值收益率完全抵消。按照之前所讨论的内容，我们可以清楚地看到，由于头寸偏好和市场趋势的作用，如果要发展一种有效的绩效基准的话，那么绩效的内容就会被排除。

从表面上看，在上升的市场趋势中，持有多头偏好的法则，是因为该法则具有预测力的实证。然而，这也并不就是理所当然的。一种法则的牛市偏好仅仅是由于它自身对多头和空头的定义而得出的。法则的多头相对于空头的条件更容易得到满足，在其他因素不变的情况下，该法则就会认为多头持续的时间要比空头的时间长。当对上涨的市场趋势的历史数据进行回溯测试的时候，这种法则就会得到一个绩效助推器。相反地，当一个法则的空头相对于多头的条件更容易被满足时，那么它会产生空头偏好，在对下跌的市场趋势的历史数据进行模拟的时候，它同样会得到一个绩效助推器。

读者也许会感到疑惑，为什么一个法则的定义可以导致多头或者空头的偏好呢？我们需要对此解释一下。回想一下二进制反转法则，即本书用于测试的类型，通常情况下要么处于多头，要么处于空头。例如，如果一

个法则的多头 (+1) 容易确定的话, 那么空头 (-1) 就很难确定。换言之, 对 -1 输出状态的寻找会受到限制, 会导致多头的持续时间比空头长。我们同样可以用公式来表示多头比空头持续的时间长。在其他因素不变的情况下, 空头出现的频率会高于多头。这就会与我们对该法则预测力的评定会受到相关限制, 或者对多 / 空头定义方式的疏忽所造成的影响的观点相悖。

分析下列的法则, 其中空头受到高度的限制, 而多头则相对宽松。该法则在标准普尔 500 指数产生的头寸是以道琼斯运输平均指数 (DJTA)¹⁰ 为基础的。假设应用于 DJTA 的移动平均线的设定范围是 (+3, -3), 那么当 DJTA 低于下端线时, 标准普尔 500 指数就处于空头, 假定这种情况属于少数, 则在大部分的时间里, 标准普尔 500 指数处于多头 (见图 1-6)。很明显, 如果标准普尔指数在回溯测试期间处于上涨的趋势, 那么该法则就会盈利。

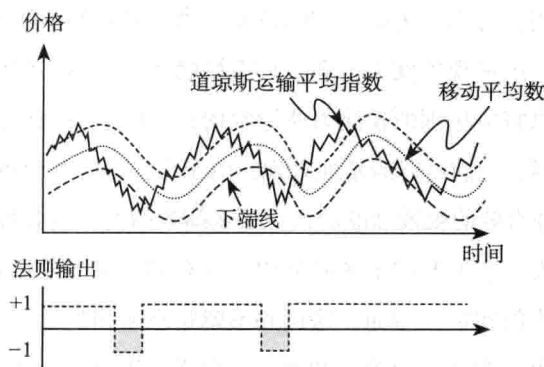


图 1-6 空头受到限制, 多头相对宽松的法则

现在, 我们来分析对两个二进制反转法则 (法则 1 和法则 2) 进行的回溯测试。测试时间是从 1976 年 1 月 1 日到 2004 年 12 月, 这期间大约 7 000 天, 标准普尔 500 指数的平均每日复合收益率是 0.035%, 或年收益率 +9.218 97%。假设法则 1 处于多头状态的时间是 90%, 法则 2 的时间是 60%, 同时, 假设两个法则都不具有预测力——就像它们的输出值是由 100 个卡槽组成的轮盘赌决定的那样。法则 1 的输出值是基于轮盘赌, 其中 90 个卡槽是 +1, 10 个卡槽是 -1。法则 2 的输出值是基于 60 个卡槽为 +1 值, 40 个卡槽为 -1 的轮盘赌。根据大数法则,¹¹ 在这 7 000 多天里, 法则 1 占

所有时间中的 90%，而法则 2 占有所有时间的 60%。尽管各种法则有不同的多/空头寸偏好，但它们具有相同的预测力——其实没有。然而，它们在这段市场历史中对收益率的预期也是不尽相同的。

一种法则的预期收益率基于三个方面：①法则处于多头的的时间比例；②法则处于空头的的时间比例（1 减去处于多头的的时间）；③在历史测试期间，市场平均每日收盘价的变动。预期收益率（ER）可以通过下面的等式来表示：

预期收益率

$$ER = [p(L) \times ADC] - [p(S) \times ADC]$$

式中， $p(L)$ 为多头比例； $p(S)$ 为空头比例； ADC 为交易市场的平均每日价格变化。

基于这种计算，法则 1 的每天预期收益率是 0.028%，即 7.31%¹² 的年收益率。法则 2 的每天预期收益率是 0.007%，即 1.78%¹³ 的年收益率。这个结果表明，法则的历史绩效在两个方面误导了我们：第一，两个法则产生的正收益率，是以它们二者都不具有预测力为基础的；第二，法则 1 的表现明显强于法则 2，即使我们知道它们都有预测力——实际上是没有的。

当对实际进行交易的法则进行测试的时候，我们可以通过以下的步骤消除由于头寸偏好与市场趋势相互作用而产生的虚假的影响。具体步骤如下：用已经观察到的收益率减去具有同样头寸偏好，但不具有预测力的法则的预期收益率。例如：假设我们不知道法则 1 和法则 2 不具有预测力，只知道它们历史的头寸偏好，法则 1 中 90% 的多头与法则 2 中 60% 的多头，而且还知道在回溯测试期间市场的平均每日收益，我们便可以计算没有预测力，但是有头寸偏好的法则的预期收益，我们已经用等式求出了已经于前文所示的预期收益。每种法则的预期收益要减去其已经观察到的绩效。因此，法则 1 的回溯测试收益率为 7.31%，我们减去 7.31%，结果为 0。这个结果可以反映出法则 1 缺少预测力。同理，法则 2 的回溯测试收益率是 1.78%，我们将 1.78% 减去，其结果也是 0，这一结果同样证明了法则 2 缺少预测力。

从本质上说：通过对不具有预测力，但又有头寸偏好的法则的预期收

益率进行回溯测试（已观察）的调整，具有欺骗性的绩效内容会被消除。换句话说，我们可以对任何法则定义基准，并以此作为不具有预测力，但是又有头寸偏好的法则的预期收益率。

解决基准的简单方法：消除市场数据的趋势

当对很多种法则进行回溯测试时，上述的过程会显得非常烦琐。每个基于特定头寸偏好的法则，都需要用单独的基准进行计算。幸运的是，我们找到了一个更简单的方法。

这个方法只需要把市场交易的历史数据（例如，标准普尔 500 指数）消除趋势，使其成为该法则测试之前的状态。我们要重点指出的是，经过消除趋势的数据仅仅用于计算该法则的每日收益率。如果交易市场的时间序列同样应用于法则的输入序列，这种方法则不能用于产生信号的预测。信号将会从实际市场数据（不经过消除趋势）中产生。

消除趋势是一种简单的转化过程，即新市场数据序列的平均每日价格的变化为 0。正如前文所指出的，如果在回溯测试中交易市场的净趋势为 0，那么该法则的头寸偏好对于绩效来说是没有任何影响的。因此，如果收益率是通过消除市场趋势的方式获得的，那么没有预测力的预期收益，或者是基准，也是 0。最后，当收益率是通过消除市场趋势的方式获得时，有预测力的法则的预期收益率将会大于 0。

为了将消除趋势的转化过程显示出来，我们首先要确定在对历史数据进行测试期间，交易市场的平均每日价格变化情况，然后再用每天的价格变化减去这个平均值。

我们所讨论的两种方法之间的等式为：①消除市场数据的趋势；②减去有头寸偏好的基准，结果可能不会一目了然。具体的数学验算请参见附录，但是，你分析这个问题的话，你就会发现，如果在进行历史测试期间，市场平均每日价格的变化是 0，那么这个法则就证明了有预测力的法则的预期收益率为 0，而不论其是否具有多 / 空头寸偏好。

为了说明这一点，让我们回到计算随机变量的预期值的公式中来，你会注意到，如果交易市场的平均每日价格变化是 0，是可以不用考虑多头 /

空头状态的。其预期收益率为 0。

$$ER=[p(\text{多头}) \times \text{每日平均收益率}] - [p(\text{空头}) \times \text{每日平均收益率}]$$

例如,如果头寸偏好是 60% 的多头和 40% 的空头,则预期收益率是 0。

$$0 = (0.60 \times 0) - (0.40 \times 0) \quad \text{头寸偏好: 60\% 多头, 40\% 空头}$$

另一方面,如果一种法则具有预测力,其在消除趋势后的数据的预期收益率将大于 0。这个正收益值反映了法则的多头和空头并不具备随机性的事实。

用每日价格比的对数形式代替百分比

到目前为止,我们一直以百分比的形式表示法则和交易市场的收益率。因为这样做有利于解释。然而,用百分比计算收益率存在着一些问题。这些问题可以通过对每日价格比,用对数的方法进行计算来消除,我们用下面的公式来表示

$$\text{Log} \left(\frac{\text{当日市价}}{\text{前一天市价}} \right)$$

用与表示百分比同样的方法是可以将以对数计算的市场收益率进行消除市场趋势的。在得出平均值之后,再用每天的数值减去这个平均值,这样就消除了在市场数据中的任何趋势。交易市场的每日价格比对数,是根据对回溯测试期间每一天的计算得出的。

其他细节:前视偏差和交易成本

俗话说,魔鬼存在于每一个细节当中。当我们对某种法则进行测试时,这个真理就会出现。目前有不只两种因素需要我们考虑到保证精确的历史测试中来。它们是①前视偏差及其相关的事件,假设的执行价格;②交易成本。

前视偏差和假设的执行价格

前视偏差,¹⁴也称为“未来信息的泄露”,即假设在某一时间点并没有

真正获得的信息我们已经知道了，信息的泄露就在进行历史测试的环境下发生了。换句话说，本来可以获得并发出信号的信息，并没有在信号假设发生的时间真正的获得。

在很多的例子中，这种问题是很敏感的。如果信息没有被认知的话，它就会夸大法则的测试绩效。例如，假设一个法则使用市场收盘价或者任何在收盘时都可以识别的输入序列。在这种情况下，对在市场的收盘价进场和离场的假设就显得不合理了。假设此举会影响前视偏差的测试结果。实际上，也许对最早进场或者离场进行假设的时机就是第2天的开盘价（假设每日信息的出现频率）。因此，该法则将在第2天的开盘价时检测假设的效果。这意味着该法则的当日（ day_0 ）收益率等于当日法则的输出值（+1 或者 -1）乘以第2天（ day_{+1} ）的开盘价与第3天（ day_{+2} ）开盘价的差。这个价格变化以对数比例的形式给出，即用第3天（ day_{+2} ）的开盘价除以第2天（ day_{+1} ）的开盘价，请看下面的公式

$$Pos_0 \times \text{Log} \left(\frac{O_{+2}}{O_{+1}} \right)$$

式中， Pos_0 = 截止到当日（ day_0 ）该法则的市场头寸； O_{+1} = 标准普尔 500 指数在第2个交易日的开盘价； O_{+2} = 标准普尔 500 指数在第3个交易日的开盘价。

上面的等式并没有显示消除趋势的法则的收益率，具体的内容请参见以下公式

$$Pos_0 \times \left(\text{Log} \left(\frac{O_{+2}}{O_{+1}} \right) - ALR \right)$$

式中， Pos_0 为截止到当日（ day_0 ）该法则的市场头寸； O_{+1} 为标准普尔 500 指数在第2个交易日的开盘价； O_{+2} 为标准普尔 500 指数在第3个交易日的开盘价； ALR 为回溯测试的平均对数收益。

当法则用滞后的，或者经过修正得出的输入数据序列时，前视偏差同样可以影响回溯测试的结果。例如，对一个使用共同基金的现金统计¹⁵法则进行回溯测试，但是这个现金统计表却延迟了2周才与公众见面，这种

情况下，我们必须把这种滞后因素考虑进去，通过这个滞后的信号反映出数据的真实有效性。本书中没有一例被测试的法则使用滞后的，或者是经过修正的信息。

交易成本

交易成本应当被考虑到对法则的回溯测试当中吗？如果这个意图的目的是在用于交易的单独基础之上运用法则，那么答案就是“是的”。例如，信号反转频繁的法则会比信号反转不频繁的法则导致更高的交易成本，这一点必须要在对两者的绩效进行比较时考虑进去。交易成本包括经纪人的佣金和股价的下跌。股价下跌是买卖双方出价的价差和投资者推动价格的下单总量，即当买入时价格上涨，而卖出时价格下跌。

然而，如果法则测试的目的是要发现包括预测信息的信号时，交易成本会使得反转频繁的法则数值变得模糊不清。因为本书进行研究的法则是以发现具有预测力的法则为目的的，而不是发现能够用作单独交易策略的法则，所以该法则是设定交易成本的。

第2章

Chapter 2

主观的技术分析的效率错觉

想法古怪的人与骗子之间的区别在于，骗子知道自己正在干着骗人的勾当，而想法古怪的人并不这样认为。

——马丁·加德纳（Martin Gardner）

本章的目的有两个。第一，鼓励对主观的技术分析的怀疑态度，因为主观的技术分析缺少认知性的内容，所以无法对其进行检验。第二，强调通过严谨的、客观的方法获得新知，对人们在缺少强有力的实证以及面对矛盾的证据时形成并持有的某些观念的倾向发出挑战。

除了我们已经接受的信仰之外，在我们大多数人的印象中，我们所持有的信念是通过对实证进行合理的推测而得出的。也就是说，我们确信某事，是因为我们对此抱有信念，我们持有某种观念是因为我们从正确的实证当中得到了正确的推理。¹ 我们知道冰激凌是凉的，地心引力是真实的存在，有些狗是咬人的，这些观点都是基于我们的第一手经验，但是在没有时间直接获取重要知识的情况下，我们更愿意接受那些我们认为可靠的二手信息。无论我们从何种途径获取知识，我们也不愿意违背自己的意愿而接受知识，或者说去相信它。

不幸的是，这种信念与我们持有的其他信念一样，全是错误的。我们就像平时呼吸空气那样，在没有经过理性思考或者可靠的实证的情况下，接受所有的观念。根据一个正处在发展之中的研究组织的报告，这种结果是由各种类型的认知错误、偏好以及错觉造成的。谎言一旦被接受，人们就会倾向于让它继续存在下去，即使出现了新的实证证明这个谎言是错误的。谎言的这种长寿性要归结为各种类型的认知错误，因为认知错误会引

导我们捍卫现存的概念。

视觉上的错觉就是错误的观点中典型的例子，因为即使事后证明它是错误的，人们依然坚持原有的观点。在图 2-1 中，线段 A 看上去比线段 B 长，但是，如果你用一把直尺实际测量一下，就会发现两条线段是一样长的。然而，这种理解并没有消除错觉。我们所看到的東西具有欺骗性，而且这种错觉会一直持续下去。

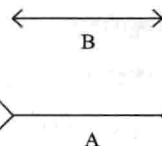


图 2-1 感觉可以欺骗我们

图 2-2² 描述了另外一种视觉上的错觉。右边桌子的平面显得更长。如果你把两幅图画进行比较，你会发现它们的面积是相同的。

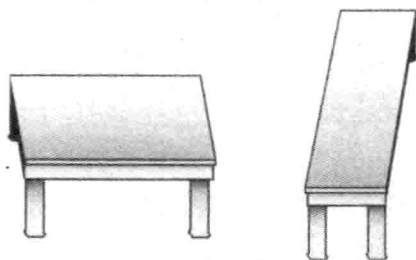


图 2-2 两张桌子的桌面在大小和形状上是相同的

资料来源：“Turning the Tables” from MIND SIGHTS by Roger N. Shepard. Copyright © 1990 by Roger N. Shepard. Reprinted by permission of Henry Holt and Company, LLC.

一般情况下，由大脑解释过的感官上的印象会产生对世界最为精确的观念，有选择的演化可以保证这种重要能力的发展。然而，与眼睛 / 大脑系统所接受的东西一样，这种有选择的进化是不完美的。在相反的情况下，由于受到形成演化的外部条件的影响，这两种系统会受到欺骗。

正因为我们会产生错觉，所以会存在知识方面的错觉。这句话看上去有点道理，实际上是错误的，这与错觉有类似之处。在超过符合认知能力演化的情况下，错误的知识才会产生。在这种相反的情况下，我们通常的学习策略就会失效，我们会发现我们知道的事实并非如此。³

在过去的 30 年里，认知心理学发现错误的知识通常是由于我们对复杂的和不确定的环境进行信息处理时所产生的系统的错误（偏差）的结果。金融市场行为是复杂而又不确定的，所以它应该对非正式的（主观的）技

术分析产生虚幻的知识这个观点并不感到吃惊。

系统的错误，有别于随机的错误，它会在相似的环境下重复发生。这其实是件好事，因为错误是可以预测的，并且通过特定的步骤可以避免其再次的发生。我们首先要搞明白的一点是，这种错误是非常普遍的。

主观的技术分析作为知识的不合理性

我们之前在第1章里对主观的技术分析进行了定义，即主观的技术分析是指自身定义不明确的、不能够以计算法则定义的，而且也不能通过计算机进行回溯测试的任何一种分析。尽管主观的技术分析范围包括很多不同的方法，但是它们却共同分享着这样一种信仰，即关于金融市场的有效的知识是可以通过非正式和非科学的方式表现的。我个人认为这个观点是不适当的。

在导论部分我们曾经讨论过，主观的技术分析并不是错误的，因为如果我们说一种方法是错误的，则暗示着这种方法是经过检验的，并且与客观实证是相矛盾的。主观的技术分析不会受到实证的挑战，因为它是不可检验的。所以，主观的技术分析并不仅仅停留在错误的这个层面，它是没有意义的。

不能够被检验的观点一定是被人们不加怀疑就接受的坊间传闻，或者是独裁者的声明。这样一来，主观的技术分析与新世纪的医学、占星术、数字命理学，以及其他不可检验的观点很类似，因为它们当中没有一个是可以按照知识的范畴进行合理分类的。

然而，许多技术分析的从业者和实践者坚决地相信主观的方法有效，当被问及自己的分析方法时，他们中的每一个人都会毫不犹豫地声称他们的观点是合理的，并且是有实证作为支持的。我个人的主张则认为这种观点是同样的认知缺陷造成的结果，因为认知的缺陷是构成错误观念和传统技术分析所使用的非科学的方法的基础。

技术分析绝对不是唯一一个受到非科学方法负面影响的领域。医学和心理学领域在面对错误的知识时，它们所遭受的惩罚要比技术分析大得多，

它们也同样受到错误的知识的困扰。这种情况已经引起了一种被称作“循证医学”的新兴运动的关注，这种运动呼吁医生禁止以证明疗效的名义进行试验。在心理学领域也有这样一个类似的运动。有一本名为《科学审查心理健康实践》（*Scientific Review of Mental Health Practice*）⁴的杂志描述了一项任务，即消除未经证实的疗法与没有事实根据的领域。不幸的是，传统的医生和心理学家不太可能放弃那些根深蒂固的方法。陈旧的观点很难得到改变。然而，新的从业者会接受劝说，并且声称运用经过证明的方法可以使病人从中受益。

尽管本章对主观的技术分析进行了批评，但是对于错误的知识来说，客观的技术分析同样会受到批评。只不过它们出现的方式不同罢了。主观的技术分析缺少数量上的证据，客观的技术分析则饱受从大量的证据中得到错误的推理之苦。关于这个问题的解决方法，我们将在第6章中进行讨论。

个人的传闻轶事：从最初真正的技术分析信仰者到后来的怀疑者

在一本关于量化的客观实证的书谈论个人轶事，多少显得有些奇怪。然而，我提供了我最初成为一名真正的技术分析信徒时的第一手资料，后来，我违背了自己的信念，并且以科学的怀疑论者的身份重新出现，作为一名科学的怀疑论者，我相信只有在运用严格的方法时，技术分析才能够提供有价值的东西。

当我还是个十几岁的孩子时，我就开始学习技术分析了，而且我相信技术分析所讲的一切都是真实的。根据图形形态来赚钱是件神奇又合情合理的事情。那个时候，我的收入主要由在夏季修建草坪、在长滩岛的北海岸挖蛤蜊以及在当地的度假胜地开垃圾车获得。尽管我从来没有回避过干体力活儿，但是不通过做苦力赚钱的想法已然成竹于胸了。

像大多数技术分析的信徒那样，我最初的信仰来自阅读那些自我标榜为权威人士的人所表述的知识。我从来都没有想到过这种合理的研究居然

没有充分的理论依据。技术分析这个术语已经带上了科学的光环。从那以后，我发现在很多的案例当中，权威人士的观点不是基于他们的第一手研究资料，而仅仅是对早期的那些自称为专家的人表述的信息的重复。幽默作家、哲学家阿蒂默斯·沃德（Artemus Ward）曾经说过，“并不是那些我们不知道的事情让我们陷入麻烦，而是那些我们认定自己知道，而实际上却是错误的知识，让我们陷入麻烦”。

我阅读的第一本书是尼古拉斯·达瓦斯（Nicolas Darvas），一位专业的舞者所著的《我如何从股市赚了200万》（*How I Made Two Million Dollars in the Stock Market*）。^①他把自己的成功归结为使用一种名为箱体理论（Box Theory）的图形分析方法。当我第一次接触技术分析的时候，买/卖信号是从市场行为中推导出来的想法是多么的新颖、让人兴奋啊！达尔瓦斯的分析方法是基于一种阶梯状的价格变化排列，即箱体（Box）理论之上的。我很快就发现这个理论的本质就是新瓶装旧酒，也就是技术分析中的支撑与阻力带的概念。但是我的热情并未因此而消退。接下来我开始研究Chartcraft的点数图的方法，并坚持对固定的几只股票绘制图形。在我16岁的时候，我的父母送给我一本由爱德华兹（Edwards）和迈吉（Magee）合著的《股市趋势技术分析》（*Technical Analysis of Stock Trends*），^②这本书被誉为技术分析领域的权威之作。正是这本书彻底改变了我。

不久之后，我开始给我的化学老师科恩先生一些股票的消息。我最初用三次曲线的方法在股市中初战告捷，并且使得更多的教师成为我的忠实拥趸。随后我又用这种方法连战连捷。没有人，甚至是我自己都没有意识到我的这些早期胜利只不过是运气太好罢了。之前，没有一本我曾经读过的研究技术分析的书籍讨论过市场中的随机性作用，或者是在某种投资方法的绩效当中存在幸运的成分。我早期的成绩之所以看上去合情合理，只是因为与书中所讲述的奇闻轶事相吻合而已，但是这种好运并没有持续太久。这种方法的缺陷开始显现，而我那些拥趸也开始远离我。然而，我对技术分析的热情依旧高涨，因为技术分析可以对我的失败进行解释。实际

①② 该书中文版已由机械工业出版社出版。

上, 我可以看到被我误读的图形的错误所在, 这样的话, 我在下一次的操作中就不会再犯同样的错误。

在上高中的时候, 我自认为是学习科学的天才, 但是实际上我知道自己对于科学来说就是个门外汉。后来, 在大学学习科学的哲学这门课时, 我知道科学不只是一组一组的事实。首先, 也是最重要的一点, 科学是从虚幻的知识中分辨出真实的事实的方法。基于这一点, 我花费了 40 多年的时间来寻找科学的哲学和技术分析之间的关系。

在 1996 年年底至 2002 年春天的这段时间里, 我作为 SLK 公司的证券自营者所积累的经验使我对技术分析产生了怀疑。自营交易者从事的工作是最为轻松和愉快的工作。他们可以自由运作公司的资本。我的交易策略是基于我 35 年前所学的有关技术分析的知识。我在 1996 年 10 月到 2000 年 2 月的前 3 年半时间里的投资是盈利的。由于我的交易方法是主观的, 所以很难说明我所获得的盈利到底是运用技术分析的结果, 还是我一贯的牛市偏好正好与市场上涨趋势相一致。我个人认为是后者的原因, 因为从我对每个月的月收益率的分析来看, 我的收益并没有跑赢市场基准, 而是仅仅与市场的基准持平。换句话说, 我没有创造显著的正阿尔法值。⁵ 在 2000 年 3 月股市趋势掉头向下运行后的 2 年时间里, 我把之前所赚的利润全部还给了市场。在全部的 5 年半的时间里, 我的投资结果总的来说, 是毫无亮点, 甚至是惨淡的。

在加入 SLK 之前, 我一直是客观交易方法的支持者。因此在 SLK 任职的时候, 我努力地研究一种系统性的交易程序, 以求提升我的业绩。然而, 因为时间有限, 而是资本不断扩充, 这些努力始终没有得到任何结果。因此, 我继续依赖经典的柱状图形的方法, 再辅以我自己进行主观解释的各种指标, 进行交易。

然而, 在很多方面我都是主张客观分析的。在我供职 SLK 的早期交易生涯中, 我开始记录详细的日志。以前的每一笔交易我都会注明所使用的技术分析理论。另外, 我的每一笔交易都是以客观的错误预测为基础的——对于这一点我不得不承认是我错了。我的经理也持有这种观点。我继续对每天的交易进行纪录, 并持续了 5 年之久。在每次交易完成时我都

会对交易进行重新分析。这项工作使我无法对错误进行合理的解释。我必须正视这些事实。正是这个原因加速了我对主观的技术分析的质疑。

由于这个结果只是对我个人而言的，我在考虑自己是否没能正确地应用我已经研究超过 30 年的技术分析指标。而且，我也开始考虑那些技术分析的观点是否真实正确了。

当我与我的同事讨论技术分析的可靠性时，他们的反映是，技术分析是完全有效的。历史图形中的形态和趋势是如此明显以至于不可能会出现错误的判断。他们认为一种被广泛应用的技术分析⁶肯定是有效的，并且会得到强有力的支持。对此，我并不同意。当我发现被技术分析认为是明显的形态与趋势，⁷可以同样有规律地出现在纯粹的随机数据中时，我对图形分析的信仰开始出现了本质的动摇。而且，研究显示，专业的图形阅读者不能够从随机过程中产生的图形中分辨出实际的市场图形，⁸这一点引起了我的注意。从实际的价格变化样本中随机绘制产生的图形，是无法预测真正的趋势和形态的。但是对于经验丰富的图形分析师来说，他们看上去就是真实的价格图形。在这个基础上，基于真实的图形所做的预测是无法让人信服的。显然，“明显的有效性”并不能够成为确定图形形态的有效性的衡量标准。

我在托马斯·吉洛维奇⁹的《理性犯的错》与迈克尔·谢尔梅尔¹⁰的《为什么人们总是相信荒诞的事情》两本书的启发下，意识到技术分析必须用怀疑的方法和严谨的科学方法来对待。

大脑：自然形态的发现者

我们倾向于寻找并且发现具有预测力的形态。由于人性受到不可预测的与无法解释的现象的限制，所以我们的认知手段会倾向于接受指令、形态以及我们从世界中所接受的感官刺激，而不在乎这些东西是否是正确的、有意义的指令和形态。“在许多案例的研究中，我们所经历的无非都是些正在发生的偶然的行爲。”¹¹“例如，月亮表面出现人脸的形状，当倒着播放摇滚乐的时候，你会收到来自撒旦（魔鬼）传递的信息，或者在医院大门

的木头纹理中看到耶稣基督，这些现象都是在出现随机的视觉感官刺激的时候，大脑做出的具有欺骗性的指令的样本。”¹²

我们具有接受逐渐形成的指令的倾向，这对于我们人类的存在是至关重要的。¹³最早的人类繁衍了很多的后裔，我们就是这些后裔的后代。不幸的是，人类的进化并没有赋予我们从无效的形态中分辨出有效的事物的能力。结果，一些有效的、有确定根据的知识就出现了错误。

不论是在史前，还是在当今的时代，获取知识可以通过两种途径：学习错误的知识，或者将其抛弃寻求真理。在我们非常容易犯的两种错误当中，人们更倾向于接受错误的知识，而忽略了重要的真理。持有进化论观点的生物学家进行了这样一种假设，即相对于早期的人类抛弃错误的知识来学习重要的东西来说，他们更愿意接受错误的思想。在狩猎之前跳舞的仪式会增加狩猎的成功概率，这个理念就是用最小的代价学习错误的知识（只要跳舞就好了），但是他们却忽略了在狩猎的过程中面对处于上风的野兽时，这种错误会导致付出生命的代价。

结果，我们的大脑就会朝着对任何事物都过度贪婪的方向进化，虽然没有做出具体的区分，但是人们更喜欢寻找有预测力的形态，以及有因果关系的东西。随着对这种方法使用的积累，错误的信念就会成为真理。有的时候，一次成功的狩猎确实是在进行完跳舞仪式后实现的。因此，迷信的实践就会得到加强。同样，当我们看到一个可以控制结果的错觉时，这种仪式就会减少人们的焦虑，即使这个结果仅仅是偶然的。

现代文明的生活极大地改变了这种观点。如今，不仅对事情进行决策变得越来越复杂，而且人们学习错误知识的成本也在不断变大。设想一下对全球变暖的争论吧。气温的升高确实暗示了一种长期的危险吗？还是说这仅仅是错误的警报呢？如果确信这个错误的警报在将来被证明是错误的，那么未来的几代人将会为此付出高昂的代价。然而，治理全球变暖的代价也许会控制的比较适度，当然前提是这个观点是错误的。

人类智力的进化是一个缓慢的过程。我们的智力能力经过了数百万年的发展，其发展主要发生在我们的祖先对环境的适应阶段。在这种情况下，大脑对此的指令是少数且清晰的：生存与再创造。人类的智力能力和

思维策略只适应于当时的环境，而不适应于 10 000 ~ 15 000 年之后复杂的现代文明。换句话说，99% 的人类进化发生在不是非常复杂的环境之下，而不是我们今天这个复杂的社会之中。如果说我们的智力错误地接受了复杂的判断，并且对我们今天所面对的事物做出了结论的话，那么也没有什么可吃惊的。我们的迷信程度与那些相信跳舞仪式的死者看上去毫无区别。

荒诞的信仰的传播

对知识既贪婪又不加选择的偏好将不可避免地导致离奇古怪的信仰。我们对错误的信仰的敏感程度在 1999 年由盖洛普 (Gallup) 对 1 236 个美国成年人进行的统计中得到了验证。人们对包括超自然等所有荒谬的说法在内的信仰的比例是惊人的。¹⁴ 谢尔梅尔公布了以下数据。¹⁵

占星术 52%

超感知觉 46%

女巫 19%

外星人着陆地球 22%

相信消失的亚特兰蒂斯大陆的存在 33%

人类与恐龙生活在同一时代 41%

与死者灵魂的交流 42%

鬼魂 35%

相信通灵的人 67%

天文学家卡尔·萨根 (Carl Sagan)¹⁶ 对信仰占星术的人比相信进化论的人多这一事实表示痛惜，他把这些不合理的、迷信的事情归结为人类当中科学文盲的比例太高，这个比例高达 95%。“在这样的氛围下，伪科学繁荣发展。奇怪的信仰，比如说新世纪的健康实践，其在某种程度上被称为科学，而实际证明它与科学没有任何关系。事实上我们并没有对此提供充足的证据，而关于其他方面的论点也被我们忽视了。然而，伪科学提出

的思想能够得以存活下来并且繁荣发展，是因为它们强有力地说出了人们在情感上的需要，即科学经常是不完整的。”¹⁷正如萨根所说，受到质疑总归不是件好事。把人们真正的需求放在一边不顾，反而去相信那些好玩儿的、只有安慰作用的事情，比如牙仙和圣诞老人，这种态度着实让人烦恼不已。

认知心理学：启发、偏好和错觉

认知心理学关注的是我们如何处理信息、描绘结果以及做出决策。它研究的是哪种感官输入可以被转换、排除、推敲、存储以及覆盖的心理过程。¹⁸

在过去的30年里，认知心理学一直在对不可信任的知识来源进行研究。好消息是对常识和经验进行的直觉解释，在大部分情况下是正确的。坏消息是人类的智力在某些不确定的情况下所做出的精确的判断是错误的，直观的判断和以非正式途径获取的知识通常情况下是错误的。由于金融市场存在高度的不确定性，错误的知识便会因此而萌发。

丹尼尔·卡尼曼（Daniel Kahneman）、保罗·斯洛维奇（Paul Slovic）以及阿莫斯·特沃斯基（Amos Tversky）所做的早期研究显示，虚幻的知识¹⁹是通过两条途径获得的。第一，人们经常受到各种各样的认知偏差和错觉的困扰，这些因素不仅会歪曲我们所经历的事情，而且还会歪曲我们对经验的理解。第二，为了弥补大脑在处理信息能力方面的限制，人类的智力经过进化形成了众多的思维捷径，我们称为判断式启发。这些在我们的意识关注之下自动发挥作用的思维法则则是以直观的判断和评估的概率为基础的。思维法则是与日常生活紧密相连的人类智力的重要标志。尽管这种便捷的思维法则总体来说是成功的，但是在某种特定的环境下，它会导致我们做出有偏差的决策，并由此获得错误的知识。

错误的知识具有不可预测性的特点，因为它本身带有自动恢复的功能。研究显示，一旦某种信仰被接受，它不仅可以经得住与自身相悖的新的实证的攻击，而且还可以经受对形成信念的原始实证的完全质疑。

接下来的这部分内容将会就认知错误对错误的知识的多样性反映进行讨论。我们会对每一种反映进行单独讨论。然而，在现实生活中，它们之间相互作用着，由此形成了主观技术分析有效性的错觉（见图 2-3）。

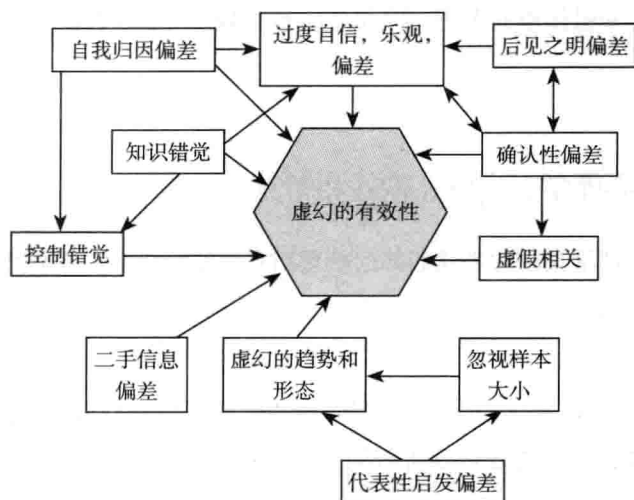


图 2-3 虚幻的有效性

人类处理信息的局限性

尽管人类的大脑有十分强大的功能，但是大脑处理信息的能力是有局限性的。诺贝尔奖得主赫伯特·西蒙（Herbert Simon）将其称为有限理性理论。“相对于需要解决的问题，人类大脑描述和解决复杂问题的能力是非常有限的。”²⁰ 因此，我们只能注意到整个世界里庞大信息量中的一小部分，并且用最简单的方法来处理。

在既定的时间内，人类大脑能够有意识地记忆单独信息模块的数量大约为 7 个，上下浮动量为 2。²¹ 当大脑需要运用结构性思维思考问题的时候，受到的局限会更大。结构性思维问题需要把大量的因素（变量）作为一个整体完整地考虑进去。研究显示，当我们必须以结构性的思维方式²² 对事情进行评估时，我们的大脑只能处理最多 3 个因素。医学诊断就是最典型的结构性思维问题。如果把一系列的 symptom 与实验室的结果结合在一起来的话，那么就可以把一种疾病从另外一种疾病中分辨出来，但是如果把两

者分开来考虑的话,则我们看到的只不过是些没有联系的、不包含任何信息的事实而已。

相对于惯性思维和线性思维,结构性思维是一种有着更高需求的思维模式。在惯性/线性思维问题中,即使变量的数量很大,相关的变量也可以进行独立的分析。每一个变量的信息是不受其他变量所传递的信息的影响。因此,一旦每个变量可以通过自身进行独立的解释,那么这一系列单独的信息就会在线性思维中相结合(例如,用代数相加法²³),从而推导出经过整理的意义。设想一下,一个序贯问题包括7个变量,每一个变量的值可以假设为+1或者-1,然后,假设5个变量的数值是+1,2个变量的数值是-1。那么线性组合或者相加性组合²⁴就等于+3,即 $(+5) + (-2)$ 。研究显示,当专家用主观的形式做出了一个多因素决定,他们主要依靠的是线性组合法则,²⁵ 尽管此举的有效性不及正规的线性回归模型。²⁶ 这些研究显示了人类的权威专家在这方面所做的不如线性回归模型有效,这是因为他们把信息和正规数学模型的一致性相结合。

当一个问题需要结构性的解决方案时,用满足序贯思维问题需求的线性方式进行的信息结合,就达不到预期的效果。在结构性的问题中,相关的信息包含在两个变量之间的关系(互相作用)网络之中。当它们处于序贯/线性问题中的时候,这就意味着变量不能够以独立的状态进行评估。为了弄清楚相互作用影响的意义,想象一下,一个结构性问题仅包括3个变量A、B和C。下一步,假设每一个变量只有两个读数,高或者低(例如,变量为二进制变量)。在结构性问题中,当因素B的值低,因素C的值高时,则因素A的读数处于高位,但是当因素B的读数处于高位,因素C的读数处于低位时,那么因素A的读数则处于高位,这种现象又说明了另外的一种情况。在序贯/线性问题中,因素A的高位读数包含了与因素B和因素C的读数毫不相干的同样信息。

两种结构的差别——结构1(A—高,B—低,C—高)和结构2(A—高,B—高,C—低)通过图2-4和图2-5来表示。图中三维空间的8个单元显示了3个二进制变量可以产生8个明显不同的可能组成的结构。然而,3个二进制变量的线性组合只能呈现4个不同的值。²⁷

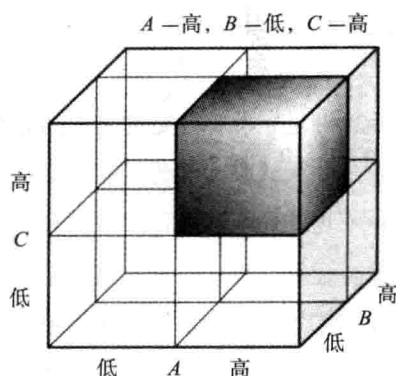


图 2-4 结构 1

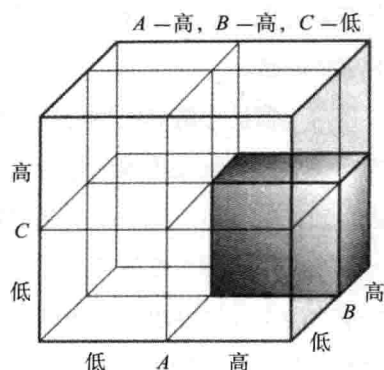


图 2-5 结构 2

主观的市场分析师试图结合 5 种指标（变量）的数值（相对中间的值）对市场做出预测，这种思维类型所面临的问题，很有可能是结构性的，而不是序贯性的。²⁸ 究竟选取合并哪 5 种指标（有些可能是冗余的，或者是根本不相关的）来做出有用的预测这件事儿本身就是一个结构性的问题。运用这 5 种指标要对可能出现的 26 种组合的预测力进行评估（10 组 2 种指标搭配的组合，10 组 3 种指标搭配的组合，5 个 4 种指标搭配的组合，1 组 5 种指标搭配的组合）。然后，一旦某一个最优组合被确认，运用多指标法则同样会引起结构性思维。就像我们前面所看到的那样，即使有的指标是二进制的，简单地说，即 3 个二进制变量可以产生 8 组不同的组合。4 个二进制变量可以产生 16 组不同的组合，5 个二进制变量可以产生 32 组不同的组合。相信自己能够做出如此业绩的主观技术分析师，可以说是过度的自信了！这并不奇怪，这种表现仅仅是过度自信偏差的一种而已。

✓ 极端的确认：过度自信偏差

一般来说，人们都太自信了。过度自信偏差是站在真理的角度上对错误进行自我评估的一种证明趋势。²⁹ 他们认为自己的个人价值和能力要比客观实证强，即便客观实证会预示着某些现象或真理。人们在很多方面都认为自己是出类拔萃的那部分人，这其中包括他们自己知道多少，是如何知道的，等等。换句话说，人们对自己掌握的知识有自大的倾向。

这种趋势是非常普遍的，以至于心理学家开始认为过度自信是人与生俱来的特征，并且对此深信不疑。据生物学家推测，过度的自信是成功交配策略的特点，所以我们有可能是那些狂妄自大的父母的后裔。高度的自信是对人性的恩惠，即使是敢于冒险的个人试图尝试新事物或者面对危险、生命无法保障的时候。在经过足够的尝试之后，有人成功了，我们也会因此而受益。另外，过度自信也会通过其他歪曲经验的认知偏差得到强化，这些偏差包括自我归因偏差、后见之明偏差以及确认性偏差，这些类型的偏差我们会在后面进行讨论。

哥伦比亚大学的心理学家珍妮特·梅特卡夫（Janet Metcalfe）用谦卑的语言对人类的过度自信做了如下概括。³⁰

人们认为他们会在不想解决问题的时候将问题解决；人们对自己有信心接近正确答案的时候过度自信，而实际上，所得出的结果是错误的；人们在还没有解决问题的时候认为自己已经把问题解决了；人们认为自己已经知道了特殊疑问的答案，而实际上他们却不知道；当没有答案的时候，人们认为答案就在嘴边；当人们没有得出准确答案的时候，他们认为自己已经得到了正确的答案；此外，他们认为自己自始至终都知道这件事；当人们认为已经精通学习材料时，实际上他们没有；当人们已经理解，甚至可以证明某些事情的时候，实际上他们仍然被蒙在鼓里。³¹

关于过度自信的一些具体发现。

（1）在处理困难或者不可能完成的任务时，过度自信是一种极端的表现。³² 研究显示，在预测赛马比赛³³的结果、从美国人的笔记当中分辨出欧洲人的笔记、预测股票价格涨跌这些事情上，都会出现极端的过度自信。在技术分析领域，也会出现类似的情况，例如，通过主观的结构性思维，结合多个指标数值对市场进行预测。

（2）过度自信与无法对各种事物进行适当的困难性评估的能力有关，而且在对越困难的事情做出评估时，这种情况就越明显。主观的技术分析师没有意识到主观的数据分析对结构性推理的需求。

(3) 随着信息碎片数量的增加,人们的自信心也随之增加,尽管预测的准确性值得商榷。³⁴ 如果独立的信息碎片不是冗余的,而是有用的,并且能够成功集成为一体的,那么,不断增加的信心就会得到确认。

(4) 自信的增加与多重输入信息的增加是一致的。当输入的信息是独立的(没有关联的),并且包含着预测信息,不断增加的信心就会得到确认,反之亦然。许多明显出现的技术分析指标是由于指标命名的惯例和对相似的市场进行具体的计算得出的。这种行为会鼓励毫无根据的信心的出现。

(5) 88%的医生相信自己在20%的病例中正确诊断出了肺炎。³⁵ 内科医生对诊断出颅骨骨折也同样充满了信心。³⁶

(6) 许多其他的专业人士对于他们自己所处的特殊领域的专业技能也充满了过度自信。例如,有90%的制药公司的经理对于自己公司的实际情况充满自信(只有10%的预期是错误的),但是他们其中只有50%的正确率。³⁷ 95%的电脑公司总经理认为自己对一般的业务行为非常熟悉。实际上,只有其中的80%是正确的。对于各自所在公司的特点,有95%的人自信非常了解,但实际上真正熟悉的人只占这其中的58%。³⁸

(7) 华尔街的分析师们对于他们预测每季度利润的能力过度自信,事实证明这是由他们获得意外收益的频率决定的。一项研究表明,预测错误的平均值是44%。这个统计数字是以50万个独立的利润预测为基础的。³⁹ 当最终的结果低于预测的范围时,人们会感到非常的意外,正是因为设定的范围过于狭小,才导致了过度自信的频繁出现。

(8) 个人投资者对于自己预测短期市场趋势和挑选股票的能力会产生过度的自信,并且认为他们挑选的这些股票的表现会优于市场,这种心理是通过他们的交易水平证明的。⁴⁰

(9) 华尔街的战略家们对一般的股票市场的预测显示出了过度的自信,但是他们也会经常吃惊于实际的结果。然而,尽管犯有这样或者那样的错误,⁴¹ 他们仍然保持着足够的信心。如果这些战略家能够真正地学会控制自己的过度自信,他们就会把大量的不确定因素考虑进去,并做出预测,而通过此举做出的预测结果就不太可能出现在设定的范围之外。但是他们

并没有这样做。

鉴于过度自信的普遍性，对于主观的技术分析的预测可能会受到过度自信的困扰。这种过度自信主要体现在3个方面：①具体方法的预测力，②主观的数据分析和作为发现新知的非正规推理的手段，③运用结构性推理对大量的指标读数和形态进行再次的市场预测的能力。金融市场的复杂性与独立的人类智力显示，对于这些领域内任何过度的自信都是毫无根据的。

在某项研究中，有两种从事预测的职业设法避免过度自信，即气象学家和预测赛马胜负的新闻记者。对于他们的预测能力的评估是经过精确测量的，因为：“他们必须每天面对类似的问题；他们做出明确的（错误的）概率预测；并且从结果中获得即时的、准确的反馈。当这些条件都不能得到满足时，对于专家与非专家来说，都需要过度的自信。”⁴²

主观的技术分析从业者在对历史进行研究的情况下，并不能够获得准确的信息反馈，这是因为他们进行预测的方法以及评估预测的准确性并没有得到明确的定义。只有用客观的方法才有可能得到反馈。然而，在实时的预测情况下，主观的技术分析从业者如果愿意做出错误的预测，那么是可以获得反馈的。错误的预测中包含认知的内容。也就是说，在做出正确预测的结果与做出错误预测的结果之间，主观的技术分析从业者建立起了清晰的、可以识别的差异。参见下面的错误预测的样本：

标准普尔 500 指数从今天开始，在未来的 12 个月当中，将会上涨得更高，并且在这段时间范围内，标准普尔 500 指数在现在的基础上的跌幅不会超过 20%。

如果市场不会在未来 12 个月出现上涨，或者说如果在未来的 12 个月当中，以当时的水平为基准的下跌幅度超过 20%。很明显，这句话中所做的预测是错误的。不幸的是，错误的预测很少由主观的技术分析从业者提出，因此，主观的技术分析从业者及其实践者也就无法获得反馈。过度自信也就得以继续存在下去。

变得伟大：乐观主义偏差

乐观主义偏差是一种延伸范围广，并且对未来充满毫无根据的希望
的过度自信。⁴³具有乐观主义偏差特点的人经常认为与其他人相比，自己的
未来会一帆风顺，基本上不会遇到不顺心的事情。

乐观主义偏差认为即使之前的预测表现得不尽如人意，技术分析师也
要坚信自己的预测技能。这一点通过《赫伯特文摘》（*Hulbert Digest*）⁴⁴ 提
供的数据得以证明，一则时事通讯通过把市场预测者的推荐转化成为一种
对价值可以进行跟踪的投资组合形式，以此对预测者所做推荐的业绩进行
跟踪。有些时事通讯栏目的撰稿人试图声明赫伯特的评估方法是有缺陷的
方式来为他们惨淡的业绩进行辩解。而在赫伯特看来，他认为那些针对他
的假设的声明是经过调整的，他需要把新闻撰稿人的模糊的推荐转变为具
体的推荐，而这些推荐者所带来的业绩是可以经过检测的。赫伯特对很多
时事通讯的撰稿人的业绩跟踪最终得出了明确的建议，即他们所做的这些
假设是根本没有必要的。

尽管他们长期的业绩表现惨淡，但是时事通讯的撰稿人仍然对自己的
预测充满信心，而且读者也相信这些撰稿人的话。这两部分人都具有乐观
主义偏差的特征。假设读者是比较现实的人，他们就会取消这些时事通讯
的订阅费，这样一来，恐怕这些撰稿人就会失业了。很显然，是某些存在
于预测力之外的因素使得这种观念得以继续存在下去。有一种可能的解释
就是撰稿人和读者都相信某种传说可以解释为什么这种基本的方法可以得
到运用，即使在实践中证明这种方法根本就不好用。随着解释的深入，引
人入胜的故事就会取代统计实证，因为人们的大脑会对这些故事越来越
感兴趣，而不是对那些抽象的事实感兴趣（参见二手信息偏差和叙述的
力量）。

自我归因偏差：让错误合理化

其他类型的偏差可以使人们过度自信的倾向加大，并且歪曲我们从经
验中所学到的东西。自我归因偏差是指以一种偏差的、自我吹嘘的方式解

释过去的成功与失败的倾向。“对不同的环境进行的大量研究显示，人们把积极正面的结果归结为自己的能力，而把消极负面的结果归结为外部的环境。”⁴⁵ 因此，我们对自己的决策和能力留下了错误的乐观主义评价，然而，过去失败的反馈会促进改变，并改善业绩，那么这些自我吹嘘的解释就会选择忽视这类的知识。这也可以解释时事通讯的撰稿人在面对被认为是消极负面的业绩时，是如何保持自信的了。

除了让我们感觉非常不错之外，自我吹嘘的解释会产生直觉。当我们努力做事并取得成功时，因果关系便会努力把这一切作为好的结果，这一点是可以理解的，而且我们还会把这种好的结果戴上真理的光环。将成功归结于好运或者其他超出我们控制范围的因素，不仅缺少情感上的感召力，而且也缺乏合理性。然而，经过我们的不懈努力最终还是失败的事情，我们更容易把这一切归结为运气太差。常识通常会忽略在好的结果中运气（随机性）成分的作用，而这些正是统计学家的工作。

对赌徒进行的自我归因偏差的研究给了人们以启发。⁴⁶ 用偏差的方式对过去的失败进行评估，他们能够在面对不断增加的亏损时保持对自己能力的信任。令人吃惊的是，过去的亏损并没有被忽视。实际上，赌徒们对待亏损投入了大量的知识活动。因此，他们得以用更加有力的角度来修正自己。亏损会被归结为运气太差。盈利，从另外一个角度来讲，要归功于天生的赌博技能。⁴⁷ 最后，不论是盈利还是亏损，都会停止向赌徒在技术上的灵感输送养分。

作为交易商，我经常听到人们把自我的吹嘘合理化。在我早期作为交易商的职业生涯中，我一直用以前已经定义好的观点坚持做交易记录。但是，后来我不得不承认我早期的预测是错误的。我的评估标准是事先经过客观定义的，尽管我的预测是以主观的技术分析为基础的。在大量的消极负面的反馈动摇我的信念之前，我还没有对我自己的能力提出过质疑。

当然，我知道这些！知识错觉

知识错觉是对我们所知道的事情产生的错误的信心，包括数量上和质量上的。它是以错误的假设为基础的，这种假设认为更多的信息应当转

化为更多的知识。⁴⁸ 当预测赛马胜负的新闻记者得到更多的信息时，他们会对自己的预测充满信心，但是他们实际的预测准确度却没有得到任何改善。⁴⁹ 额外的事实产生了更多的知识错觉。

知识错觉与技术分析关系密切，因为在大量的数据面前，知识错觉最容易出现。不仅金融市场产生独立且连续不断的时间序列，而且还可以衍生出更多的时间序列，我们将之称为技术信号。这一点足以使技术分析的从业者认为自己拥有更多的知识，但实际上并非如此。

知识错觉在某种程度上来源于大脑对结构性推理和大量的变量做出同步解释的过度自信。实际上，独立的大脑在进行结构性推理（见第二部分）时，只能处理 2～3 个变量。由于没有认识到这些局限性，人们很容易认为越多的变量（指标）会导致更多的知识以及更多合理的观点。

我可以解决它：控制的错觉

控制错觉是指我们高估了自己对事件的控制程度。那些觉得能够控制自己的人要比那些认为不能控制自己的人更高兴、更放松。⁵⁰ 这些认识上的歪曲是由自我归因偏差引起的，而与之相反的是过度乐观主义偏差。

根据诺夫辛格（Nofsinger）的说法，最容易引起控制错觉的 5 种行为具有以下 5 种特点。⁵¹

- （1）高度的个人参与；
- （2）大范围的选择；
- （3）可以获得用于思考的大量信息（知识错觉）；
- （4）对工作的高度熟悉；
- （5）在活动中取得的早期成功。

前 4 种特点可以应用于主观的技术分析。高度的个人参与是以对市场数据频繁的分析、对新的指标用新的方法来解释、反复绘制趋势线以及反复计算艾略特波浪等为基础的。大范围的选择是指市场朝哪个方向发展，用哪种指标，在什么地方画趋势线等。进行市场分析时所使用的方法与平时研究使用的方法非常相似。第 5 个因素，早期的成功，是偶然的。有些人会获得初始的成功，是因为自我归因偏差把成功归结为人们的专业

技能与技术分析方法的有效性，而不是归结为偶然的因素。所有这些因素可以引起并保持着让人们自己拥有赚取跑赢市场收益率的能力和对未经证明的事物的控制感。

后见之明偏差：我知道事情会朝着那个方向发展

后见之明偏差会产生这样一种错觉，即人们觉得在对某一事件进行回顾，并且已经知道结果的情况下，对未知事物的预测要比对实际存在的事物进行预测容易得多。一旦我们知道了未知事件的结果，比如说哪支球队会赢得比赛的胜利，或者说出现某种技术分析形态后，价格会向哪个方向运行。在这种情况下，我们倾向于忽视在我们知道结果之前，我们对结果是多么的不确定（见图 2-6）。

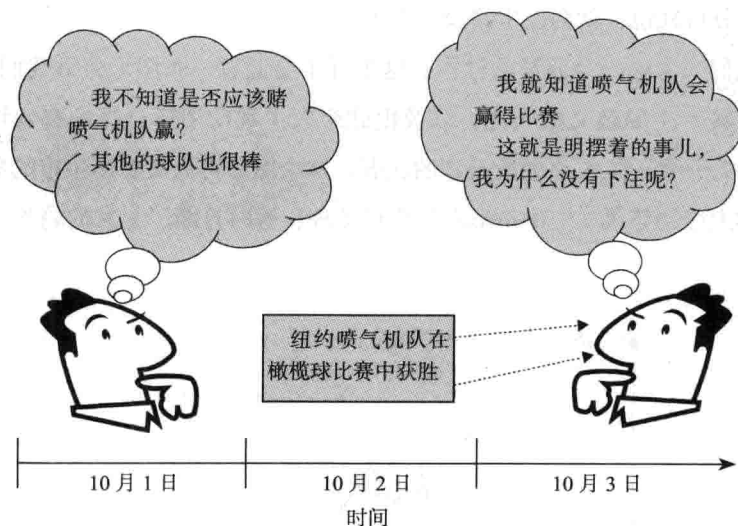


图 2-6 后见之明偏差

这种认知上的歪曲可以通过大脑对记忆的储存这种方法来理解。过去的事情被储存在相似事件的类别中（例如，运动事件），而不是作为连续的事件被记录。正是出于这种原因，发生在不同时间的相关类型的事件就会在此时混合在一起。结果，在一场橄榄球赛的结果出来之后，赛前的不确定状态就会与赛后明确的结果相混合。赛前种种模糊的迹象的这一本质是导致我们出现不确定性的主要原因。你可以在赛前预测任何一支球队将获

胜。赛前的状态与赛后的状态的混合是即时的，且无意识的，这使得我们无法回忆起我们在赛前的想法与质疑的状态。结果，赛前不明确的迹象在比赛结果出来之后就显得不是那么明确了。这种状态就产生了错误的自信，使得我们相信自己有预测的能力。由于这种缺陷的存在，科学家们对于进行预测并检验评估其准确性的定义过程就显得格外关心了。

主观的技术分析容易受到后见之明偏差的困扰，因为它缺少对模式识别、预测生成以及预测评估进行明确的定义。在现实中，主观的技术分析从业者面对着一项模棱两可的任务。一个简单的图形里包含着很多种可能的预测的暗示，范围从极端的熊市到极端的牛市之间。然而，当同样的图形被拿出来回顾时，分析师再对此进行预测就会显得模棱两可，其预测力也会随之下降，究其原因，是因为我们已经知道结果了。这对于主观的图形技术分析来说，就会产生错误的有效性。

我们来考虑一个假设的样本，这个样本会通过一张图形告诉我们结果是如何将产生的歧义最小化，以及由此夸大主观的图形分析的有效性的。如图 2-7 所示，在 A 点处所看到的状态。这个图形可以被看作牛市的象征：前期的上涨趋势再加上牛市形态的旗形整理，暂时打断了上涨的趋势。

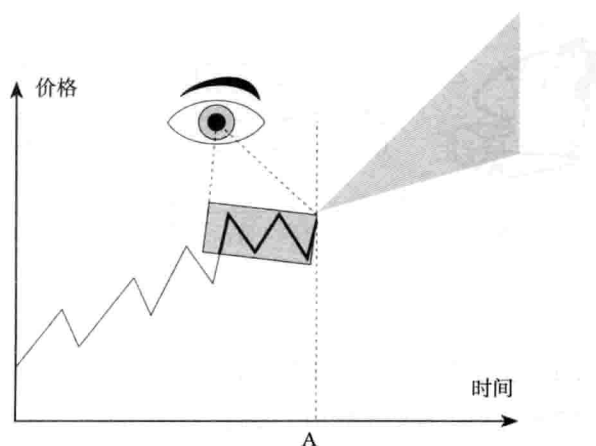


图 2-7 牛市判断

同样的图形也可以理解为熊市走势：头肩形态再加上跌破颈线位，然后在反弹的过程中在某一理想的时间上做卖空交易（见图 2-8）。除了价格

的涨跌以外，请注意两个图形都止于 A 点。我们所做的预测全部基于此图形形态。

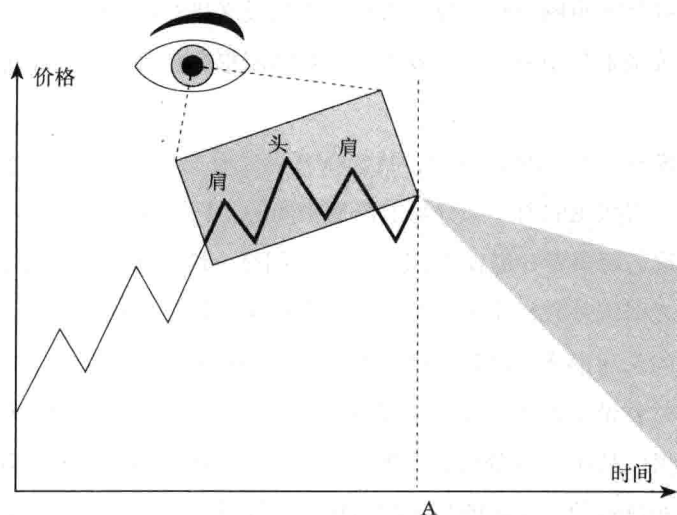


图 2-8 熊市判断

实际上，在图形中的 A 点处具有支撑牛市或者熊市预测的技术特点，并且是产生歧义的焦点之所在，因为 A 点从技术特点的分析来看，是既可以支持牛市预测，又可以支持熊市预测的，如图 2-9 所示。我们会解释技术分析师们在面对同样的历史图形往往做出不同预测的原因。

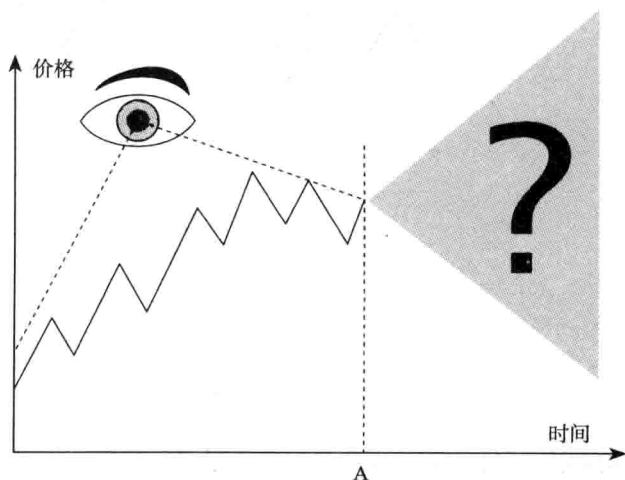


图 2-9 预测的真实不确定性

现在我们来想象一下图形分析师在以后的时间看到同样的图形（B点）时的表现。我们在A点之后会看到两种不同的走势结果。此举将会显示在A点时，后见之明偏差是如何无法对产生的歧义进行区分的，所以，不论实际上的价格走势朝哪个方向发展，主观的图形分析师都会产生虚幻的有效性。

首先来看一下在图2-10中，我们在B点处看完整整个图形后，技术分析师的反应。这张图的前半部分与上图是相同的，只是在A点之后添加了新的内容。这名观察者知道在A点与B点之间是上涨趋势。依我所见，我认为后见之明偏差鼓励主观的分析师注意到牛市特征的旗形形态，而不是熊市特征的头肩形态。换句话说，当我们看到结果时，会倾向于在绘制精确预测的时候做出并不显著的错误预测。实际上具有歧义的形态对于观察者来说，仅仅是在A点处所包含的信息，即对于从B点往回看的观察者来说，图形所显示的是明显的牛市旗形形态。同理，从B点处进行观察的观察者就不会看到显著的熊市头肩形态。因此，在B点所留下的痕迹证明了图形分析的有效性。

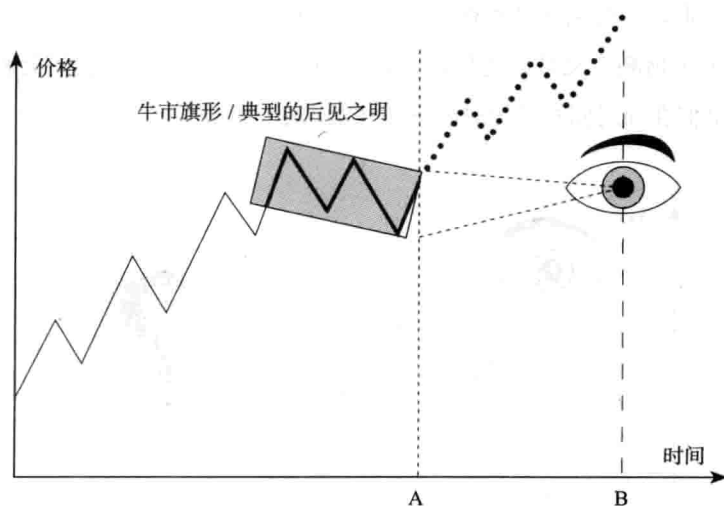


图2-10 后见之明的虚假确定性 1

如果从时间点A到时间点B的价格走势是下跌趋势，如图2-11所示，我认为后见之明偏差会导致观察者在B点准确地看到头肩形态，但是观察

者是不会注意到前面已经证明是错误的牛市旗形形态。

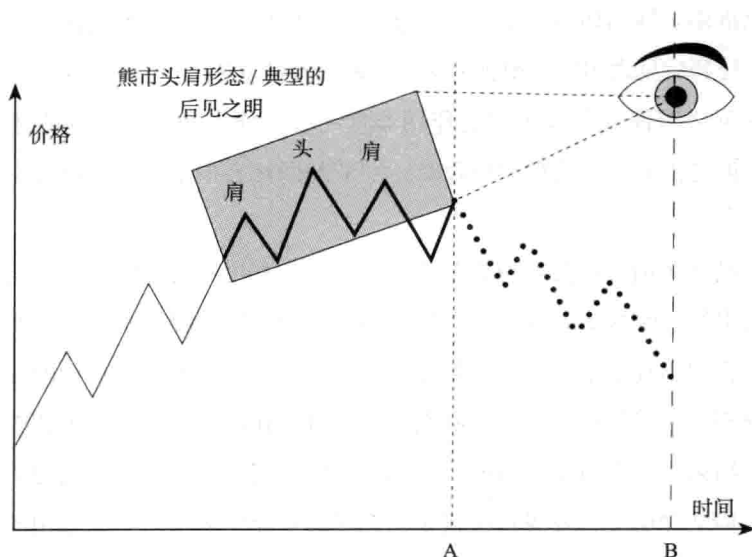


图 2-11 后见之明的虚假确定性 2

这些演示是人为虚构出来的。不论结果如何，我们在一个牛市、熊市信号相互混合的综合体中，通常会选择那些有准确预测力的特点的图形。在主观的图形形态已经确定的情况下，这种模棱两可的方法缺少客观的评估基准，对于主观的技术图形分析师来说，运用后见之明偏差的结果将会产生虚幻的有效性。

我们中的有些人是第一次遇到头肩形态，他们并没有研读历史的图形来证实头肩形态是否与其他人所描述的相一致。对于我们中间没有留下“这个方法真管用”印象的人来说，是不是我们就是后见之明偏差的牺牲者呢？

我们通过一例有说服力的实证实验⁵²来证明人们遭受后见之明偏差的困扰。在一项特别的研究中，学生们被问及尼克松总统 1972 年访华可能出现的结果的概率。例如，实验要求学生预测尼克松总统是否会与毛泽东主席进行会晤，以及美国是否在不能满足中国所期望得到的外交承认的情况下与中国建立外交关系。这两个事件在尼克松出访之前都是不确定的。当两个星期之后事情的进展得到公布时，实验要求学生回忆他们之前所

做的预测。在经过两周的间隔之后，67% 的学生记得他们的预测要比实际的情况准确。换句话说，学生们不能回忆起他们之前的不确定性。几个月以后，受到后见之明偏差困扰的学生比例已经从 67% 上升到了 84%。

后见之明偏差在庭审时的作用非常明显。证人相信自己的陈述准确无误，但是他们对事情顺序以及具体细节的回忆会根据事情的进展而出现改变。⁵³

后见之明偏差也会对历史叙述造成影响。利用后见之明的历史学家通常会指出第三帝国的兴起是可以预测的。他们声称纳粹的种子已经明显地在第三帝国出现之前就在各种著作中初露端倪。实际上，第三帝国的兴起只是众多可能猜测中的一个，我们完全有可能在阅读时就把这种猜测加入其中。结果，对于历史学家而言，第三帝国的兴起看似是不可避免的，但是当我们回顾历史，发现这仅仅是在无尽的可能中被历史学家采用的一种可能罢了。

其他的实证实验显示针对于降低后见之明偏差的策略收效甚微。⁵⁴ 即使在人们被告知警惕后见之明偏差并设法避免时，后见之明偏差仍然会发生。它的出现超出了理性的控制范围。甚至专业的技术都不是那么有用。在一项研究当中，实验要求一组医生评估其他医生的诊断错误。医生们进行的评估是基于最终确定的病理报告，并且用与疾病相关的知识把自己武装起来。评估者无法理解这样的错误会出自一位经过专业训练的内科医生之手。（这一次，结果再次使得过去的预测看起来要比现实更加具有预测性。）

认知过程对后见之明偏差的反应是怎样的呢？尽管事情还没有得到澄清，但是我们已经认为自己非常聪明，并且可以控制一切了。一种得到了支持的推测提出，这与记忆的工作方式有关。或者说，记忆根本就没有工作。⁵⁵ 用简单的术语来表达，就是大脑对记忆的记录不是按照顺序进行存储的，并且随时以原来的顺序进行重演的被动过程。相反，记忆可以将事件解构之后再进行储存，并且在需要的时候把这些事件重新搭建起来。对于富有经验的事件，记忆可以把它们切分并储存在相关的类别里，而不是按照发生的时间顺序储存。“上个月在乡下的朋友家遇到了一条狗”这句话

可以被切分开来存储在明确的类别之中，并且进行神经定位；对于狗、朋友、到乡下旅行的记忆，也许发生在过去不同的时间段内。当我们试图回忆起上个月去乡下旅行时，大脑会把储存在单独类别里的信息重新集合，并且重新构建记忆。心理学家推测我们的大脑就是通过这种方式进化的，因为我们的大脑是有效并且高效的。然而，由于事件发生的时间并不是储存系统的重要特征，所以人们很难回忆起事件的发生顺序。它体现的只是事件的顺序。例如，我们什么时候接受的知识。当我们评估自己的预测能力时，这一点是非常关键的。

在将事件进行解构储存的过程中，新的经验会与旧的经验相融合。就像一滴墨水掉进一杯净水当中一样，水与墨汁会不可避免地完全混合在一起。新获得的知识，例如，某件不确定的事件的结果，就与之前获得的结果相混合。后获得的知识与之前获得的知识彼此之间无法明确区分，并且看上去好像自始至终都是已知的。所以我们就不难理解人们很难重新建立之前的不确定性的原因了。

主观的技术分析师可以克服后见之明偏差吗？在回答这个问题之前，我们必须把分析师在两种不同的任务中的情况考虑进去：①寻找形态——运用历史数据的预测力寻找形态；②进行实时的预测——用当前时间的形态进行新的预测。研究形态的工作是以道氏理论的创始者，查尔斯 H. 道 (Charles H. Dow)，或者艾略特波浪原理的鼻祖——拉尔夫 N. 艾略特 (Ralph N. Elliot) 为代表的。他们在形成自己思想的过程中，非正式地提出或者检验了各种假说，例如哪种图形形态具有预测力。运用当前时刻的形态做出实时的预测，这项工作是以运用道氏理论的道氏理论家——理查德·罗素 (Richard Russell) 和艾略特波浪理论的专家罗伯特·普里切特 (Robert Prechter) 为代表的。

在对历史进行研究的情况下，我主张后见之明偏差是不可避免的，这是因为它不可能回避结果知识。仅仅知道价格的走势，会使得分析师在对任何正在进行评估的方法是否具有预测力的理解造成偏差。只有主观的技术方法能够提供避免后见之明偏差的机会，因为只有某一时点上所获得的信息才能够发出信号，并且用客观的方法对这个信号进行检验。

在对当前的时刻进行预测的时候，主观的技术分析师可以避免出现后见之明偏差，但是前提是他们愿意做出错误的预测。如果在做出预测的时候，分析师提出的这个结果构成了预测错误，或者用于检验的程序与将要应用的程序一样时，这个预测就被视为是错误的。例如：

(1) 市场将从现在开始连续 6 个月走高，并且在这个时间框架内，市场不会从现在的水平下跌超过 20%。

(2) 或者，市场从现在的水平下跌 20% 之前，将会在现有的水平基础之上上涨 20%。

(3) 买入信号出现，在卖出信号出现之前，会持续多久。

前两种预测结果很明显是错误的。例如，如果市场在上涨 20% 之前下跌 20%。这个预测就是错误的。仅此而已！第三个预测，一个信号清楚地暗示了评估的过程——即从买入日期直到卖出信号出现时这段时间的盈利与亏损。错误的预测为分析师及他们的客户提供了宝贵的反馈。遗憾的是，很少有主观的预测者这么做。

二手信息偏差：好故事的力量

包括专家在内，没有一个人可以通过直接经验获得所有必需的知识。因此，我们从那些自称什么都懂的人那里学到二手知识就是不可避免的了。

在传递知识的无数种方法中，叙述或者讲故事是目前最为流行的方式。生物学家史蒂芬·杰伊·古尔德 (Stephen Jay Gould) 把人类称为“会讲故事的动物”。我们千百年来一直在以这种方法来分享思想。结果，我们所了解的大部分事物中，包括如何算术，都是以叙述的形式储存在大脑里的。⁵⁶

由于这个原因，好的故事要比客观事实有更强的说服力。心理学家推测具体的、有色彩的以及感人有趣的故事会对我们的信念产生极大的影响。这是因为这种故事可以让我们认为它是早已存在的故事。⁵⁷ 简单和抽象的概念不足以点亮大脑中的故事网络，因为只有令人兴奋的故事才能做得到。

哲学家罗素曾经说过，当我们通过非正式的（非科学的）方式学习知识时，我们往往受到“情感与兴趣的影响，并非是知识的数量”。⁵⁸为此，科学家们训练他们自己用一种完全相反的方法将这些知识进行重新推演，也就是说减少对这些故事的关注，并着重于客观的事实，当然，如果能够减少知识的数量，那就更好了。

并不是所有的故事都是如此的扣人心弦。为了吸引注意力，一个故事必须有趣并且让人易懂，而且不能因为是已知的东西而让我们感到反感。我们大多数人都会被对真实的人进行的描述所吸引，尤其是那些我们知道的人。让我们感动的故事需要有令人愉快且有益的特点。好的传说不仅可以使我们兴奋，而且还要反复证明自己存在的价值，并让它传承下去。

真理与传说的冲突

我们求知的欲望与通过好的故事获得的知识之间存在着冲突。幽默作家 H. L. 门肯（Menchen）曾经雄辩地说：“围绕在真理周围的主要是令人不舒服且无聊的东西。人类的大脑所寻找的是更加有趣、亲切的东西。”但是现实却并非如此，所以听众必须依靠说故事的人实事求是的讲述，而不是让他们说那些我们愿意听的话。

令人遗憾的是，即使二手信息的主要意图是以讲述吸引人的故事来传递知识，但是它也会经常为了满足人们的兴趣而出现偏差。说故事的人深知讲述充满修饰的故事是十分乏味的。当紧密相连的部分被无限放大时，之前的那些不一致的、模棱两可的事情就会显得不那么重要了。⁵⁹ 尽管经过修正的观点叙述起来会更容易被人们所接受，但是它要表达的真理就有可能因此消失了。最后，听众会无限夸大这些信息的准确性与有效性。即使科学研究的论文对这些信息进行概括和宣传，但是这种情况还是会发生的。在无数次的重复之后，真理就会被人们进一步抛在身后。起初作为一个有意义的结果可能最终被报告为具有更多意义的发现。

特别具有说服力的叙述往往是那些与因果关系紧密相连的事情联系在一起。因果关系的概念是一种自然的认知能力，⁶⁰ 因果链可以满足我们对发生在身边的事情进行解释的需要。⁶¹ 对陪审团的决定所进行的研究显示，

在法庭上所做得最好的因果叙述会赢得官司。⁶² 将证据以最佳的方式变成一致且令人可以接受的故事，在今天是非常流行的。⁶³

谬误在被视为合情合理的情况下是很难被察觉到的，这就是对因果进行解释时所面临的实际问题。有一个自从 20 世纪 50 年代以来就广为流传的故事：十戒原则包括 297 个字，独立宣言包括 300 个字，林肯的葛底斯堡的演说包括 266 个字，但是一份来自政府价格稳定办公室对于甘蓝菜的价格管制的指令却有 26 911 个字。而实际的真相是价格稳定办公室根本就没有发出过这样一份指令。尽管如此，这份指令仍然吸引了人们的注意力，即使政府机构出面向大众澄清这件事是完全不存在的。这件事一直到现在还在被人们谈起。经过修饰的指令很容易就被描述成为“联邦政府的指令”。⁶⁴ 这一事件的反讽特征与冗长的官僚主义在人们心中的合理存在使得这件事不会就此简单地从人们的视野中消失。

艾略特的传说

引人注目的故事也许能够解释艾略特波浪原理（Elliott Wave Principle, EWP）——技术分析中的一种夸大的推测方法，具有持续的吸引力的原因。艾略特波浪原理认为价格波浪可以反映普遍规律以及事物形成的方式，这一规律不仅可以体现在金融市场的波动中，而且贯穿于整个自然界，例如贝壳、银河系、松果、向日葵和其他自然现象。

根据艾略特波浪原理（EWP），市场趋势是不规则的，即每一个波浪当中包含着其他与之相同的波浪，其范围可以从仅仅持续数分钟的微波到持续上千年的宏波。⁶⁵ 这种共享的形式，就是艾略特波浪原理（EWP）。它包括涨跌在内的 8 个阶段的价格运动。实际上这种涨跌形态的普遍性不仅适用于金融市场价格的变化，而且也在大众心理学、文明的兴衰、文化形式以及其他社会趋势的演变过程中得到证实。这一理论声称它几乎对所有经历过发展和变化循环的事物进行了描述。即使是艾略特波浪原理的主要支持者罗伯特·普雷切特的交易生涯，也会遵循一系列的涨跌原则来证明艾略特波浪原理（EWP）是正确无误的。这一理论与斐波纳契数列以及黄金比例等关于数字排列顺序的理论联系在一起。⁶⁶ 这些理论都具有数学的性质。

然而，那些声称可以解释任何事情的理论实际上什么也没有解释。作为使用最为广泛的理论，艾略特波浪原理（EWP）并不具有合理性，不过是一个如罗伯特·普雷切特⁶⁷所讲的引人入胜的故事而已。但是艾略特的理论也是有一定的说服力的，因为市场历史中最小范围的价格波动能够与该理论中所述的8个波浪中的任何一个阶段相吻合，所以艾略特波浪原理看上去还是非常有效的。我个人认为这种现象是由于该理论松散的定义规则以及对大量嵌套波浪的变化幅度所做的推测造成的。即使在知道地心说的理论是错误的情况下，哥白尼之前的天文学家依然可以对进行观察的行星的运行做出解释。艾略特的分析师现在就是这么做的。中世纪的天文学家把天文学中本轮之间套着本轮的说法与他们所掌握的数据相吻合，这一点与艾略特波浪原理（EWP）的分析师把一个波浪中包含着其他若干波浪的理论与市场数据做到一致是类似的。因此，如果艾略特波浪原理（EWP）的基础概念是错误的，分析法也能够做到与过去的数据相一致。⁶⁸

实际上，任何灵活的模型都可以完美地与以前的一组观察相吻合。例如，包括一些术语的多项式函数（数量与数据点的数量相同）也是可以对数据进行完美改造的。然而，具有可以无限的做到与过去的数据相一致的能力的模型或者方法，如果不能够对未来做出可以检验的（错误的⁶⁹）预测，则这个模型或者方法就应当被认为是既没有意义的，也没有用处的。

艾略特波浪原理（EWP）也可以应用于天文学，尽管该理论认为银河系是具有对数螺旋性质的。但是其在这方面的表现却不尽如人意。⁷⁰那么接下来如何解释这个理论持久的吸引力呢？我认为要把这归结为这样一个事实，即艾略特波浪原理（EWP）体现着这样一种因果关系，它承诺对市场的过去进行解释，并且对未来做出预测，在这方面这种方法比其他任何一种技术分析方法都要好。正如我们在前面所讲的那样，有些传说因为太过引人入胜，而不能让它就这样消失。

根据个人的私利编造的故事

个人的一己私利可以加速对二手信息的扭曲。人们的意识形态和理论观点使得他们倾向于有选择地将故事的某些方面刻意尖锐化，并将其他方

面最小化，使之与他们自身的观点相一致。具体的技术分析方法或者技术分析的宣传者普遍上都具有清晰的意识形态和对金融的兴趣。

这种指责可能指向了我、⁷¹ 或者任何正在争辩这个观点的作者。然而，当讲故事的人受到客观实证和重复程序的约束时，对于私利所产生的扭曲的影响来说，就没有任何可以发挥的余地了。

✓ 确认性偏差：现存的信念是如何过滤经验的，以及矛盾的实证是如何存活下来的

“一旦某种观念形成，我们将以支持这个观念存活下去的方式对信息进行过滤。”⁷² 确认性偏差把我们以前确认过的信念视为可信赖的新的实证，而把与之相反的信念视为不可信赖的实证。⁷³ 这种倾向阻碍了我们从新的经验中学习、并且消除错误的思想的机会。确认性偏差对我们坚持错误的信念的原因做出了解释。

确认性偏差的理性基础

在许多实例当中，对能够支持现有信念的实证给予特殊的照顾，对与之相悖的实证提出质疑是可以理解的。如果没有这种规则的存在，信念能否得以存活下来是存在相当大的不确定性的。

只要以前的知识能够被有力地支持着，那么将以前对知识的运用比作过滤器是人类智慧的特点。科学家们对一则声称核聚变是在一台从当地的五金商店里购买的零件所组装的仪器，并在室温的温度下完成的说法提出了合理的质疑。⁷⁴ 科学家们的怀疑最后通过一项客观的测试得以证明，即他们不能复制所谓的冷聚变。几乎没有人会认为超市张贴的小报上所写的“猫王会乘坐 UFO 回归地球，并且建造一个外星人主题公园”的信息是真实的。然而，以前的信念之所以不能够被证明，是因为它们缺少了从确凿的实证当中获得有力的支持，因此，授予以前的信念智慧的守护者的地位是不理智的。问题是人们并不在意他们所持有的信念是无法证明的、不合理的。最后，本不该发生的确认性偏差依然会出现。

偏差的认知

确认性偏差是感知能力起作用的结果。个人的信念产生对事物的预期。反之，则通过某种感知形成结论。因此，我们就看到了我们想要看到的东西，并且理所当然地得到了我们想要的结果。就像亨利·大卫·梭罗所说的：“我们所听到的和所理解的只是我们一知半解的东西。”常言道，眼见为实，这远比眼不见就相信的说法更有说服力。

我们通过下面的实验来说明期望对认知所造成的影响。当实验对象接到一杯人们认为含有酒精的饮料（实际上不包含酒精）时，他们减少了焦虑感。然而，当其他的实验对象被告知将会得到一杯非酒精饮料（实际上是酒精饮料）时，他并没有因此而减少焦虑。⁷⁵

我们不但没有意识到新的信息与以前的观点的冲突，而且也很难接受与我们已经习惯的观念相悖的结论。这样做的后果就是当我们对与以前存在一致的观点进行严格的检验，并且剔除错误的信息时，我们更倾向于接受与过去已经存在的观点相一致的信息。确认性偏差解释了我们不能对新信息做出改变的原因。⁷⁶

更糟糕的是，当确认性偏差与随机定律相结合时，随着时间的推移以及人们对这种观点的熟知程度的加深，做出错误的观念的预测的可能性也大大地增加了。⁷⁷没有预测力的技术分析（以掷硬币作为信号）仅仅能够依靠运气获得偶尔的成功。随着时间的推移，这些经过实证的实例就会逐渐积累。相对于失效的方法，确认性偏差的重要性就更加明显了。自欺欺人的恶性循环就是最好的证明结果。这一点暗示了随着时间的推移，对于有缺陷性的方法的有效性所持有的信念就会得到加强。因为太多的基于机会和运气获得的成功会不断地暴露于大庭广众之下，而它实际的优势却被人们所忽视了。所以，那些曾经使用过的无效的技术分析的人就会意识到这些方法的缺陷。

动机因素

动机的因素也驱动着确认性偏差。技术分析从业者在进行感情和金融投资时，都有各自认为最好的方法。这一点对于将职业生活与特别的方法

相结合使用的从业者来说，是再确切不过的了。

还有一种与人们自身的系统性观念和态度保持一致的强烈动机。由费斯廷格（Festinger）⁷⁸ 阐述的认知失调理论认为人们会很积极地去减少或者避免心理上的矛盾。⁷⁹ 人们对与所持有观点不一致、相矛盾的证据所引起的不安，使得他们很难接受这些实证。

偏差问题与搜索

确认性偏差同样倾向于寻找新的证据。搜索偏差会增加遇到新的证据来证明原有观点的机会，从而降低遇到非确认性或者矛盾的事实的可能性。这种情况通常发生在进行主观的技术分析研究的情况下。根据定义，搜索偏差局限于寻找有支持性的传闻证据。因为从主观上定义的模式并没有明确预测错误的条件，所以矛盾的例子很难被发现。

受过科学的方法训练的人在面对这种倾向时，要做两件事。首先，在测试的开始，为评估的结果设定客观的基准。其次，努力捕捉与以前的观点和假设相矛盾的实证。能够有力地支持寻找反驳证据的想法要比仅仅通过确定的证据进行证实的想法更具可信性。

客观的技术分析也会产生确认性偏差。假设一个分析师正在用回溯测试法寻找可以盈利的法则。如果经过测试，第一条法则的业绩并不能够让人满意，那么就会继续寻找更好的方法。这种方法将会对一系列包含可变参数、逻辑、指标的规则进行测试，直到找到最佳的方法。我们将这种方法称为数据挖掘。由于何时结束这一程序完全由分析师所决定，所以寻找的范围也就不会受到限制。这样就保证了有着良好的过往业绩表现的法则被发现了。研究者最初对于搜索好的法则的信念便由此得到确认了。

实际上，由数据挖掘发现的某一法则过去的业绩表现夸大了其在未来的表现。这种夸大未来的表现称为数据挖掘偏差。我们在第 6 章将会具体讨论。这种问题的出现并不是数据挖掘本身的问题。实际上，当各方面都处理妥当时，数据挖掘就是一种非常有创造性的方法。因此，错误就没有被考虑到由于数据挖掘所导致的上涨偏差中去。考虑到搜索的程度可以导

致发现好的方法，对于一种法则未来的盈利潜力做出明确的推理是完全有可能的。

模糊的评估基准对确认性偏差的影响

由于主观的技术分析在如何对其自身的预测进行评估这一点上的表述是含糊其辞的，所以，对其预测成功或者失败的实证证据也是不明确的。这种不确定性提升了分析师们指出以往已经运用的预测的能力，同时也避免了对预测误差的确认，进而有效地使主观的技术分析从业者免受消极反馈信息的影响，从而使得改变他们信仰的事情发生。⁸⁰

主观的技术分析方法的预测误差在某些方面是让人很难理解的。其中之一就是对形态的重新命名。无效的向上突破形态将会被重新定义为成功的牛市陷阱。根据技术分析的教义，向上的突破预示着上升的趋势（见图 2-12）。

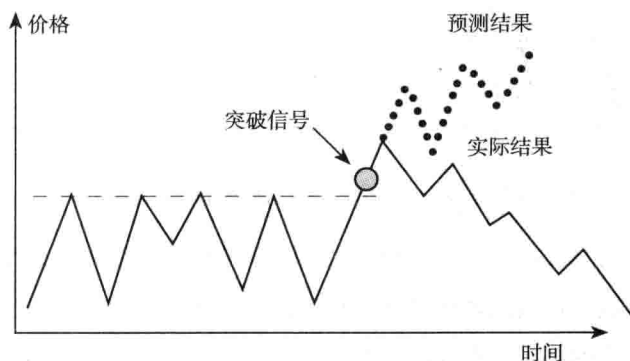


图 2-12 无效突破

当一波上涨的趋势没有最终实现时，这种情况就应当视为预测误差。如果可以替代的话，这种形态可以称为牛市陷阱，⁸¹ 在图 2-13 中，我们在另外一种形态中发现了明确释放的信号。客观的形态定义和浅析的重新定义评估基准避免了对实证证据进行改动。

当预测缺少一个明确的结束点时，另外一种让人感到费解的预测误差就会显现。预测的结束点是对将来的预测进行评估的时间或者事件。具有明确的或者暗示着一个明确结束点的预测或者信号，实际上是毫无意义的。

技术分析的权威人士通常发出没有明确结束点的预测：“我对股市中期的走势持看涨的观点。”这样一个模糊的声明为预测者提供了足够的后续空间来避免自己的错误。不论是已经确定的时间期间，还是具体的后续的事件。比如，相反价格运动的固定的百分比，避免了预测者参与到对证据进行后续的、秩序混乱的行动中来。而这些预测都是通过好的方法描述出来的。然而，“我希望市场下跌 $X\%$ 之前，在现有的水平上至少上涨 5% ”这句话就没有为造假留有余地。客观信号方法表明了一个明确的结束点，即对现有的头寸（持仓）给出了进行下一步操作的信号。

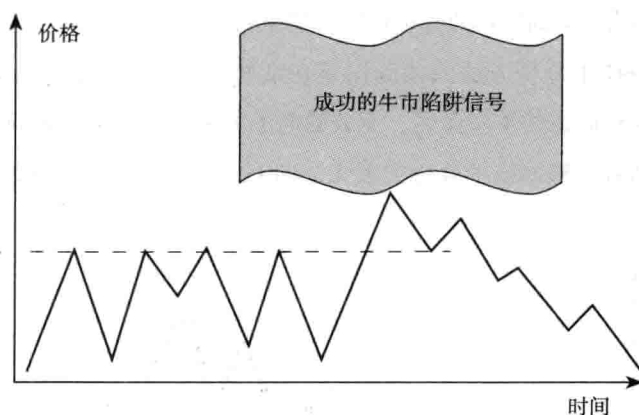


图 2-13 重新命名的形态导致的模糊的预测错误

直到事实可以被认为具有支持性的证据时，才存在改变预测评估基准的可能。这一点很清楚地解释了科学的两面性。当然，科学所具有的灵活性使得它能够接受所有能够进行检验的思想。然而严格的测试检验与评估还是要进行的。实际上大多数的科学家对于一个毫无意义的思想的检验，还是通过严格的程序来进行最终判断的。如果人们对每天的生活方式都是按照相同的模式进行的，那么他们就会更倾向于接受错误的观点。⁸² 人们善于编造想法、理论来解释貌似合理的经验。⁸³ 然而，人们不太善于检验那些自己已经形成的观点。其中最主要的原因之一就是人们无法明确定义哪种结果符合具有支持性证据的标准，而哪条又不是。在没有客观评估基础的情况下，为我们已经相信的事物寻找具有支持性证据的机会与动力就不会受到约束，并且朝着不好的方向发展。

确认性偏差与模糊实证

确认性偏差对以前的观点进行证明的能力主要依靠新实证的真实程度。当新的实证并不明确的时候，确认性偏差操作起来就会不受约束。因为模糊的实证既可以被解释为是矛盾的，又可以被证明是得到支持的，通常情况下，模糊实证会因为得到支持而被人们所接受。

主观的技术分析方法的模糊性有两个方面，即形态与信号是如何被定义的，以及基于这两者的基础上所做的预测是如何进行评估的。因此我们需要找到能够支持这个的实证。不充分的实证不足以改变已经深入人心的观点，但是对于不充分的实证的解释存在很大的不确定性，所以某些解释也许会与原来已经存在的观点相一致。不充分的实证的典型特征就是对于它自身的解释可以适用于各种不同的观点。例如，某些价格形态图形与艾略特的脉冲式图形，或者说完全有别于主观的形态，或者是纯粹的股价随机运动⁸⁴相吻合。而客观上定义的形态和评估标准就不会出现在这种情况下，无论数据是否作为某一特殊形态的样本，也不论结果是否与对形态的预测相一致。

确认性偏差与矛盾实证

我们现在来考虑当一个实证明显有悖于原有的观点时，确认性偏差是如何表现的。在大多数情况下，实证就是一个混合体。有些实证明确支持原有的观点，而有些则与之矛盾明显。即使在这种情况下，确认性偏差仍然起着作用，只不过其表现方式更为复杂了。

也许有人认为自相矛盾的实证会被忽视，或者说一反常态地成了具有支持性的实证。但是，这种趋势是不会发生的，因为人们把自己比作理性与认知一致的结合体。“人们为了看到他们想要看到的，以及相信他们希望相信的事物，不愿意忽视那些矛盾的实证。”⁸⁵然而，人们为了降低自相矛盾的实证与原有的被人们广为接受的观点之间的冲突，通过采取改变矛盾的信息这种方法来实现他们的目的。换句话说，确认性偏差鼓励人们对明显自相矛盾的实证进行巧妙的认知修改，来减少/降低它的重要性，以及降低与确认性实证和人们所珍视的信仰之间的冲突。

人们通过应用更加严格的接受标准来抵制自相矛盾的实证。⁸⁶ 对处于有利位置的实证所使用的标准，是要求这些实证具有合理的特性或者存在确实有效的可能性。反过来，对于与原有的受人们所珍视的观点相悖的实证所使用的标准，则要求它具有**超越任何可能的质疑的说服力**。对于信仰疗法的信徒来说，一个不存在疑问的确凿的说法是可以接受的。然而，受约束的严谨的科学研究拒绝承认信仰疗法的有效性，它会对每一个可能的证据，不论合理与否，提出质疑。由于对矛盾的实证的要求十分苛刻，所以它可以非常有说服力地面对任何质疑，而那些所谓的和谐的实证只具有微弱的一致性，这就使得错误观点的持有者能够继续坚持他们的观点。

对确认性偏差进行研究所得出的一项出乎意料的发现，就是与原有观点相悖的实证实实际上会加强对原有观点的影响。人们普遍认为具有辩证性的实证可以减少对原有观点的影响力。然而研究所显示的却是相反的一面。例如，在一项实验中，受试者要对赞成或者反对死刑表达各自明确且具有说服力的证据。⁸⁷ 受试者分为两组，一组赞成死刑，而另一组反对死刑。两组受试者要对死刑进行的不同研究做出总结和评论。其中一项研究表明死刑对犯罪具有有效的威慑力，而另一项研究则表明死刑对犯罪根本就不起作用。⁸⁸ 这项实验显示了受试者更倾向于接受对自己原来持有的观点起到支撑作用的实证，并认为与原有观点相悖的研究是有缺陷的和证据不充分的。当对有关死刑的观点的实证公布出来的时候，受试者感到他们对原有观点的认同度又得到了加强。换句话说，一个具有辩证性特征的实证会使参加实验的两组受试者对各自原来持有的观点更加偏执。

另外，考虑到受试者在研究中所没有做的事情也是非常有益的。不为人们所接受的实证既不会受到误解而成为受大众欢迎的证据，也不会被人们所忽视。相反地，它还会对现有的缺陷进行仔细的检查，以降低实证本身对于原有观点的影响。换句话说，就是对有重要意义的认知工作进行解释，或者将自相矛盾的实证最小化。⁸⁹

一项与技术分析相似的实验分析了确认性偏差是如何让一个失败的赌徒在赌博亏损越来越大的情况下，能够保持更强的翻本儿的幻想的。输钱表面上是无效证据作用的结果，然而，赌徒却对此非常乐观。托马斯·吉

洛维奇指出，赌徒的所作所为是通过他们对赢钱和输钱用带有偏见的方式进行评估所得出的。⁹⁰假设在上文中参与实验的受试人所得出的结论是真实的，那么赌徒是不可能忘记或者忽视矛盾的实证的（失败的）。相反地，他们会把失败转化为贬义的证据。这些在赌徒的日记中可以找到。赌徒们把更多的精力用于关注输钱，而不是获得的盈利，他们对盈利和输钱的解释不尽相同。他们将盈利归功于自己对赌博的敏感程度（自我归因偏差），而把输钱归咎于运气太差，如果对投注策略稍加调整，那么赢钱的就是他们了。

以上的这些发现显示了主观的技术分析从业者在面对挑战和时下流行的方法相悖的客观实证时所做出的反应。⁹¹例如，假设把一种主观的分析方法转换成为客观的分析方法，然后测试/证明这种方法是无效的。那么这种方法就会致力于发现大量的证据来评估这项研究。任何科学领域的研究都会受到各种理由的质疑。主观的技术分析从业者也许会炫耀捡樱桃的例子来证明他们使用的方法的成功。最后，当人们抱怨象牙塔里的专家们把精力用在与本职工作毫不相关的地方的时候，主观的技术分析从业者很有可能比任何时候都要坚定地放弃他们所使用的分析法的价值。

超越确认性偏差：真正的信徒

我们已经看到错误的观点是以模糊的概念为基础的，并且有可能在需要对需要让其成为非主流观点的要求时，反而变得更加发展壮大。然而，一种观点有着强大的生命力。研究表明，一旦某种观点生根发芽，它就可以完全不受能够成为一种观点的实证的质疑。

研究显示，当受试者被误导形成错误的观点时，他们所持有的观点会一直持续到他们被告知受骗为止。一项研究之所以与技术分析有关，是因为它包含了对形态的识别。受试者被要求从虚构的自杀信息中分辨出真实的自杀信息。⁹²在开始的时候，受试者会收到伪造的反馈数据，以此来让他们相信自己已经学会了如何从虚构的信息中分辨真实的信息。而实际上，受试者根本没有学会如何去做。随后，实验者将真相告诉了受试者。虽然证实了提供的证据是虚构的，但是在某种程度上，受试者仍然相信他们可

以从虚构的自杀信息中分辨出真实的信息。实验者在其他类似的实验中也发现了相同的效果，即一个观点在经历对其全盘否定之后，仍然还可以存在。⁹³

这些发现预测了如果拉尔夫·艾略特、W. D. 江恩和查尔斯·道这些人如果还在世，并且宣布他们的方法就是一场有预谋的骗局时，那么他们的信徒将做何反应。也许现在的技术分析实践者还会继续相信这一点。假设艾略特波浪原理分析，或者说它的某种说法，现在已经简化成为了一种客观的运算法则，那么当客观检验失败的时候，我们反倒是对艾略特波浪原理的信徒们的反应饶有兴趣。⁹⁴

主观分析方法最有可能影响确认性偏差

一些主观的分析方法是鼓励确认性偏差的。这些主观的分析方法具有以下3个特点：①对某种方法为什么能够有效运用做出复杂的因果性解释，或者编造出令人信服的故事；②高超的改造能力，即准确配合或者解释市场行为的能力；③没有产生错误的（可检测的）预测的能力。

符合以上特点的技术分析方法包括：艾略特波浪原理、赫斯特周期理论（Hurst cycle theory）、占星术预测以及 W. D. 江恩分析和其他分析方法。例如，艾略特波浪原理就是基于复杂的因果解释来证明在物质世界、大众心理、文化和社会中形成的普通力量。而且，它具有高超的改造能力，即应用大量可以改变持续期限和幅度的嵌套波浪，导出任何以前的历史数据片段的艾略特波浪的总数。然而，除了艾略特波浪原理其中的一个说法外，⁹⁵艾略特波浪原理这种分析方法是不可产生可检验的/可以被检验的预测。

我们首先来考虑第1种因素，即复杂的因果解释。基于令人难以理解的因果故事是可以经受住实证的挑战的，因为它们所说的使人们感到有必要去了解这个世界。我们必须对自身的经验做出解释，而且事后我们也具有将刚刚发生的事情编成看上去貌似合理的故事的能力。“为了让编造的故事可以维持下去，那么就要对不同的结果、特征以及原因进行解释、证明和寻求一致性。我们通过实践学习，用更快、更有效的方法应对这个课

题。”⁹⁶ 在一项研究中，人们总是受到鼓动而形成错误的观点，然后又要寻找实证来为这种观点构建理论基础，那么这些人就要比那些没有被告知、解释的受试者更容易抵抗这种观点所带来的影响。⁹⁷ 那些拥有充足的理由去解释自己观点的受试者开始变得更加自信了，即使当他们被告知他们之前是被人利用而接受错误的观点之后，他们还是选择继续相信自己的观点。⁹⁸ 实际上，他们与其他没有被告知所持观点是受人利用的结果的受试者一样，对这种观点深信不疑。

第2种因素可以使技术分析对实证的挑战产生免疫力。它们具有把对过去市场运动的改造与对未来市场运动无法做出可检验的（貌似合理的）预测相结合的能力。利用技术方法把过去所有的市场运动趋于一致的能力，可以起到清除任何矛盾的实证的效果。当结果与事先的预测完全不符时，错误的预测便发生了。即使是主观的评估也是无法隐藏错误的。此时，标准的解释就是错误的预测是由错误的应用分析方法造成的，而不是分析法自身的缺陷所造成的。而且还要运用分析法来确定这个明显不对的结果对这种谬论的支持。能够获得良好追溯的能力与无法进行预测的能力相结合，可以消除与技术分析方法相悖的实证出现的可能性。

在科学的领域中，淘汰错误理论的基础理论实际上是现有的错误理论对未来进行的预测，它与实际的预测是相矛盾的。对科学家来讲，这就是这种理论需要淘汰或者进行重新定义的信号。如果二者之间是矛盾的，那么新的理论就会对一组未来的预测进行检验。这种淘汰虚假理论的基本技术将在图2-14中显示，并且在第3章中做进一步的讨论。

关于是否修正或者拒绝一种技术分析方法或者其他类似的技术分析方法，来自可以证明此分析方法所做的错误预测的明确实证。如果随便采用事后修正的方式来消除错

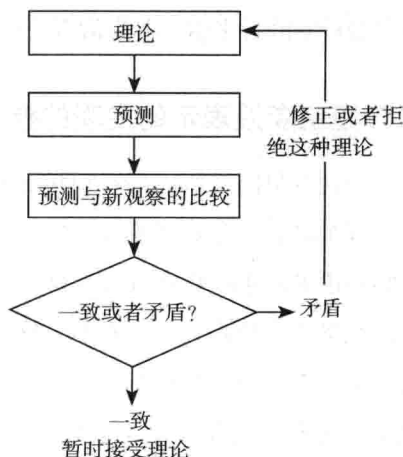


图2-14 科学是如何淘汰错误理论的

误，那么拒绝分析方法的理论基础就毫无意义了。

其他研究提出人们更加倾向于确认性偏差。奇怪的是，这些人都是知识分子。哈佛大学心理学家大卫·珀金斯（David Perkins）安排了一次十分有趣的研究，研究表明人越聪明（以 IQ 值为基准），就越能够为他们的理论观点构建清楚的理论基础，并保护它们。⁹⁹ 然而，珀金斯同时发现这些人有很强的能力对现有的观点进行理论基础的构建，所以，这些高智商的人就很少考虑备选的解释与观点。因此，他们更容易形成观点，而不受其他观点的影响。社会科学家杰伊·斯图尔特·斯内尔森（Jay Stuart Snelson）称其为思想免疫力。他声称（宣布）高智商与成功人士很少改变他们预先的假设。¹⁰⁰ 这种想法在从事特殊领域、积累重要知识的人群中的效果最为显著。由此我们也想到那些高智商又有经验的主观的技术分析实践者也是可以拒绝他们之前深信不疑的某些分析方法的。

虚幻的相关性

对于虚幻的相关性的研究，讨论的是人们对自己以及自身能力的错误认识问题。现在我们要研究的是人们对世界的认识，即一组变量之间是否存在预测性的关系。¹⁰¹ 研究显示人们非常倾向接受事物间虚幻的联系。虚幻的联系是指对一组变量之间的关系产生了错误的感觉。这点对于主观的技术分析实践者来说是非常值得关注的问题。

以二进制变量表示的主观的技术分析法

一组变量间之所以是互相联系的，是因为其中的一个变量相对于另外一个变量来说，是具有预测力的，它可以预测另外一个变量，也就是我们通常所说的结果变量或者因变量。尽管主观的分析法并不经常使用这些术语进行描述，但是几乎所有的分析法都被视为在一组变量中确实是有关联的存在。技术分析法中所说的具有预测性的事物是指诸如头肩形态这类的技术分析形态或者信号。这类形态和信号是以预测在此图形出现之后价格的走势为目的的。当头肩顶形态的图形形成之后，据此做出的趋势判断应

该是下跌趋势。实际上，主观与客观的技术分析的运用（活动计划），也被视为对这种联系的探索。尽管这些方法在大多数情况下会通过组合的形式预测市场的趋势，但有的时候也是独立工作的。

对主观图形形态的预测和结果，同样可以理解为是二进制变量。也就是说，假设现在只有两个变量，对于预测性变量来说，图形形态与信号只会出现两种情况：①出现信号；②没有出现信号。¹⁰²而对于结果变量来讲，对未来价格走势做出预测的形态或者信号也只会出现两种结果：①确实出现信号；②没有出现信号。

例如：一般情况下对于头肩形态的预测通常会与下跌趋势相联系，当在某一点位上出现头肩形态，并且不论颈线在此处被突破的信号是否已经发出，那么伴随着信号所出现的下跌趋势则不一定会出现。¹⁰³因此，对主观的图形形态或者信号的理解就会全部基于目前的这个框架内。

本章的中心论点就是主观的分析法的知识具有虚幻的有效性的特点。这部分分析的是具体的有效性错觉，即虚幻的相关性。虚幻的相关性是指对两个毫不相干的变量之间存在相关性的错觉。

为什么我们会有虚幻的相关性倾向

大量的研究表明人们容易受到虚幻的相关性影响。¹⁰⁴这些虚幻的东西看上去来自错误的信息处理。我们往往对明显错误的事物给予太多的关注，并且对与错误相关的东西坚信不疑，而很少将注意力集中在那些既没有与以前的预测相关，但没有被确认的事物上，也没有将精力放在与现实完全矛盾的事物上。在技术分析中，确认性的实证是指那些形态/信号已经明确发出，而结果也是按照信号所预测的方向去发展的案例。例如，当头肩顶形态已经出现，下跌趋势随即展开的时候。如果你相信某些特殊的形态/信号的有效性，试想你所相信的理论是否是以完全或者主要根据已经经过实践成功检验（经过确认的案例）的例子为基础的。现在来想一想你所熟悉的类似的案例吧（错误的积极信号或者错误的消极信号）。¹⁰⁵

然而，确认性的实证仅仅是一组二进制变量所产生的4种结果之中的1种。其他3种结果是②形态/信号出现了，但是根据图形预测所出现的市

场趋势却没有发生；③形态 / 信号没有出现，但是根据图形形态预测的结果却发生了；④形态 / 信号没有出现，而根据形态 / 信号的预测应该出现的结果也没有发生。类型②和③是预测错误，类型①和④是正确的预测。我们把类型②称为错误的信号或者错误的积极信号，类型③称为错误的消极信号。

左上角单元格的诱惑

一组二进制变量所产生的 4 种可能的结果可以用规格为 2×2 的列联表(交叉表)(见图 2-15)来表示。这种普通的分析工具包括 4 个单元格，每一个单元格代表 4 种可能出现的结果中的一种，在每个单元格中输入一个观察对象所产生的其中一种结果的实证数量。假如主观因素妨碍了数据的获得，这个表格对于阐述主观的技术分析的图形形态以及结果就十分有用。

		预测结果出现	预测结果没有发生
形态 / 信号出现		正确的预测 确认性实证	错误 错误的信号
形态 / 信号没有出现		错误 无效的信号	正确的预测 没有机会

图 2-15 2×2 列联表

研究表明人们容易受到虚幻的相关性的影响，因为他们过度关注确认性实证的数量，如图 2-15 左上角单元格。¹⁰⁶ 同时，他们没有对其他 3 个单元格中案例的数量进行适当的分析。我们把造成这种错误的原因归结为两点。首先是与确认性实证的特点有关，其次是与构成对存在着的的相关性的结论给予充分实证支持的概念所做出的错误的直觉观念有关。

实证的特点指的是它如何明显、生动或者突出地表现出来。显著的实证会吸引人们的关注。具有确定性的样本是突出的，因为它们确认了受到质疑的假说。换句话说，它们是具有回报性的。当主观的技术分析专家认为他们已经发现了一种有用的图形或者信号，他们通常会分析历史数据来

确认他们最初的假设是否正确。在这种情况下，具有确定性的实证，也就是那些当图形形态出现后，并做出正确预测的样本，会引起人们高度的重视，这是因为它们验证了分析师们最初的推测。试想当 R. N. 艾略特（R. N. Elliott）在最初形成他的预感（即所有的市场运动全都符合同样的基本形态）在经历了 5 段脉冲式的波浪后，后面有 3 段修正的波浪这一理论之后（5-3 形态），接下来会发生什么呢？此时，他的假设并没有获得让人信任的地位。这不过仅仅是一种令人质疑的假设罢了。然而，一旦这个假设出现在艾略特的脑海中，它就会激发一种为该假说寻找实证的动力。换句话说，是艾略特最初对自身的验证才使得确定性的事实变得突出。

必要的实证 VS. 充足的实证

一旦确认性实证站稳脚跟，那么错误的直觉就会自然而然地形成。这种直觉是对造成相关幻觉的反应，对实证数量的关注就将会被认为是建立有效相关性所必需的条件。一般来说，实证本身就是很充足的。但这个说法是完全错误的。

这也并不是说这种直觉是错误的，甚至是不完全的。直觉明确地告诉我们确定性实证（图 2-15 左上角单元格所示）对于建立相关关系是多么的重要。如果头肩形态，或者任何形态和信号是有效的预测，那么形态以及期望的结果也会出现。然而，对于大量的确定性实证来说，足够建立相关关系的假设就是错误的。换句话说，尽管确认性实证很重要，但是它们并不是通过自身建立相关关系的。实际上，分布在 4 个单元格中的实证数量必须要考虑有效的相关性是否存在的问题。

由于我们没有意识到现有的观点并没有得到充足实证的支持，所以我们看到的不是某种意见或者信念，而是从客观实证中得到的合理推测。¹⁰⁷ 大量研究表明，¹⁰⁸ 当人们用非正式的分析法来确定两个变量间是否存在关系的时候，他们更容易犯两种错误。第一种错误是当两个变量的相关性非常强¹⁰⁹（隐形的相关性）的时候，他们没有观察到两者间的相关性是有效的。第二种错误是他们没有意识到相关性是无效的（虚幻的相关性）。这两种错误与观察者对来自以前已经存在的观念的期望密切相关。

首先，我们考虑一下当给人们展示的两个变量（实际上相关联，但并不是真的希望两者之间有联系）。¹¹⁰ 在一项实验中，我们设定相关性的范围值是 0 ~ 1，除非二者的相关性超过 0.70，否则人们是无法观察到相关性的存在的。在上面的这个数值范围内，1.0 代表完全相关，而 0 则代表不相关。在技术分析的领域中，数值为 0.70 的相关性是极有可能出现的。¹¹¹ 这一发现显示依赖于非正式 / 直觉的数据进行分析的主观的技术分析师，是容易错过有效的相关性的，而这种相关性是他们所没有想到的。

与虚幻的技术分析相关的问题是具有对不存在的相关性进行察觉的倾向。当观察者持有以前关于在图形形态与结果之间存在关系的观点，或者有一种动机发现我们之前所提到过的确认性偏差时，即使是根本不存在的相关性也会被认为存在相关性。所以我们很难想象一个技术分析师是不会受到以前的观念、某种促动因素，或者两者都包含在内的因素的影响的。

即使没有相信以前的观念，或者不存在某种促动因素，人们依然表现出对发现虚幻的相关性的倾向。斯梅斯隆（Smedslund）早期的一项实验曾经对这种倾向性做过研究。¹¹² 假设给一些护士看 100 张卡片，这 100 张卡片代表了 100 个实际的病人，每张卡片显示了病人假设被证明或者没有被证明没有特殊的症状（二进制预测），并且不管最后是否诊断成为具体的疾病（二进制结果）。护士们的工作就是确认这种症状是否与疾病有关系。护士对此做出的预测在图 2-16 中显示。

正确的诊断			
		疾病出现	疾病没有出现
症状出现	单元 A	37 案例	17 案例
	单元 B		
症状没有出现	单元 C	33 案例	13 案例
	单元 D		

图 2-16 护士估算的数据——斯梅斯隆的研究（1963）

超过 85% 的护士认为症状与疾病具有相关性。左上角的表格（症状出

现与疾病的出现)显示了他们确信的 37 个突出的确认性实例。然而,他们的结论并不能够得到数据的支持。对数据做适当的分析,也就是说把所有的 4 个单元格全部考虑在内,显示了这种症状与疾病本身确实没有关系。

给出所有预测的单元

将所有的数字全部考虑在内的正式的统计学测试,也许是最为严谨的方法,用来确定是否护士的数据是有效关联的预测。然而,我们不用采用正式的测试来证明症状与疾病之间是不存在相关性的。我们只需要计算每个单元格中的简单比例。例如:已经确认有症状,且最终发病的病入的比例大约是 0.69,即单元 A/(单元 A+单元 B), $37/54=0.685$ 。而没有发现症状,最终却发病的病人比例是 0.72,即单元 C/(单元 C+单元 D), $33/46=0.717$ 。这个简单易懂的结论显示不论症状是否存在,最终发病的可能性是基本相同的(0.69 和 0.72)。如果两者的比例差异过大,例如:有症状且最终发病的比例是 0.90,而没有症状,最终也发病的比例是 0.40,那么症状与发病之间就是确定存在相关性的。但是这并不能够说明症状就是引起病发的主要原因,只能说是对导致病发的有效预测而已。

根据简单的比例法所显示的症状与病发之间没有相关性的事实,竟然有 85% 的护士做出症状是对病发的最可靠预测,这实在是太让人吃惊了。护士同样被给予了其他的数据集。经过解释的因素是否能够被护士认为是具有相关性的,简单地说就是确认性实证的数量(单元 A:症状与疾病同时发生)。他们的判断并没有受到之前计算得出的比例的影响。换言之,护士没有对在其他 3 个单元中的案例数量给予关注。就像我们之前提到的那样,许多心理学研究都对斯梅斯隆的发现¹¹³表示支持。

发现两个变量之间相关性的正确的统计学方法是卡方检验(Chi-square test)。¹¹⁴当变量之间不存在相关性时,这种方法主要是确定分布在每一单元中的实证数量是否与对该单元预期的实证数量相一致。换句话说,这项测试是用来确定单元里的数量是否与四个单元中随机的少量实证相背离。当单元格中的数量与随机形态明显相背离时,¹¹⁵卡方检验的统计数据就会出现偏大的数值。这就是相关性的合理实证。

图 2-17 是把对护士的研究数据进行改变之后，对于某种症状与病发的相关性，单元格中的计算结果是如何得出的。每一个单元都包括如果症状与病发之间没有预测关系的话（随机形态），对实证的预期数量。例如，在单元 B 当中，症状与病发同时发生的实证（错误的信号）有 7 例。如果这些实证的例子符合随机形态，则预计会出现 16 个错误信号。因此，如果症状出现了病发的比例与症状没有出现病发的比例之间的数值比是 0.87 VS. 0.50。所占比例的不同显示了症状与疾病是有相关性的。

		最终诊断	
		疾病出现	疾病没有出现
症状出现	单元 A	47 期望值 =37.8	7 期望值 =16.2
	54		
症状没有出现	单元 C	23 期望值 =32.2	23 期望值 =13.8
	46		
	70	30	100

图 2-17 有效相关性的实证

虚幻相关性中的不对称二进制变量的作用

当一组变量中包括不对称的二进制变量这种特殊类型的二进制变量时，虚幻的相关性就会发生。¹¹⁶ 二进制变量可以是对称的，也可以是非对称的。如果二进制变量中每一个变量的大约数值与存在的属性（红 / 蓝、共和党 / 民主党），或者是发生的事情（雨 / 雪、洪水 / 火焰）相关的话，就属于对称的二进制变量。然而，在不对称的二进制变量这种情况下，变量的两个值当中的一个描述了一种属性的缺失或者并没有发生的事件。例如，（认为存在 / 认为不存在）或者（事件发生 / 事件未发生）。换言之，其中的一个变量现在所处的状态代表了现在没有发生的事情，或者是根本就不存在的某种特征。

研究显示，¹¹⁷ 当预测和结果是不对称的二进制类型时，确认性实证（形态的出现和预期的结果出现）是十分突出的，并且特别有可能消耗运用非正式数据的分析师的注意力。向左上角单元格转移的注意力有力地支持了对虚幻相关性的认知。回想一下当变量是不对称的二进制变量时，其他3个单元的变量所包含的实证的情况。在这3个单元中，每组变量的一个或者全部变量正在记录着形态的缺失、没有出现的预期结果，或者将两者同时记录。

一个没有发生的事件，或者当属性的缺失是很容易被人所忽视的，因为人们很难对实证的种类具有某种概念，并且对这种概念进行评估。研究暗示了我们需要更多的认知工作来评估那些没有发生的实证。为了说明额外的认知工作，需要对包含没有发生的事件，或者属性的缺失的陈述做出处理。让我们考虑以下两种陈述，尽管两个例子都表达了同样的信息，但是第1种表述的意思是一种肯定的形式，非常直观。然而第2种表述是以否定的形式出现的，这就需要我们用心去思考一下了。

表述 1：所有的人都会死。

表述 2：所有不会死的东西都是人类。

很多主观的技术分析论点都坚持，在一组不对称的二进制变量之间确实存在相关性的观点，即：一个信号/形态要么存在，要么不存在，以及信号/形态所做预测的结果要么出现，要么不出现。因此，早期的心理学研究提出主观的技术分析的实践者非常有可能受到虚幻相关性的欺骗。他们会对确认性实证的印象十分深刻，即信号/形态不但出现，而且预期的结果也必然会出现。一些主观形态及其结果的例子请参见表2-1。

表 2-1 主观形态及其预期结果

形态 / 信号	预期结果
头肩形态	下跌趋势
完整的艾略特 5 段波浪	正确的波浪
强势上涨，然后进行旗形整理	涨势的持续
极端的牛市行情解读	熊市

隐藏或者丢失的数据使得虚幻相关性的问题更加严重

我们到目前为止所讨论的一切显示了人们是如何因为对单元格中所有

相关的信息缺少充分的考虑，从而导致了错误的、相关性的信任。然而，主观的技术分析师所面临的困境并未因此而结束，比发现虚幻相关性的这种倾向更严重的问题就是数据丢失。列联表中单元格中的计算结果数据并不是轻易就能获取的。数据丢失问题可以说是另外一种主观分析的结果。

在没有对信号/形态，以及评估其预测结果的标准进行客观定义的情况下，是很难获取列联表中所需要的数据的。没有这些数据，信号/形态的有效性是无法进行评估的。尽管主观的技术分析方法的实践者看上去从来不缺乏确定性的样本，¹¹⁸ 我们还是要将这种方法应用于确定性实证的结果，并显示在左上角单元格中。客观的定义是解决数据丢失问题唯一的方法。也许直到技术分析接受客观现实的时候，技术分析才会在奇幻思维和神话的基础之上继续发挥作用。

一旦某种虚幻的相关性被人们接受成为事实，我们之前提到的两种认知错误就会使得这种错误的想法长期保持下去，对似是而非的联系的信任将会增加对额外具有支持性的实证的关注（确认性偏差）。而且，这种所谓的相关性会被人们的观念体系轻易地吸收。正如吉洛维奇所指出的，一旦某种现象被认为是存在的，人们就会对其做出解释来证明它存在的原因，以及它代表的含义是什么。换言之，即人们非常善于为既定事实编造各种故事。¹¹⁹ 研究表明，当人们在收到错误的信息时，例如，被告知在某些工作中的表现或好或差时，人们通常会马上寻找理由来解释。¹²⁰ 一部分主观的技术分析所提出的观点或者是使人着迷的故事，其本身就可以解释它为什么会如此吸引人，而故事本身也支撑着这种观念。

行为主义者对虚幻相关性的解释

迄今为止，我们所讨论的内容通过认知心理学的方法，对错误的观点的产生和维持进行了解释。作为心理学的独立分支学科——行为主义，在某种程度上对主观的技术分析的有效性这一错误观念的持久性做出了不同的解释。

行为主义心理学关注的是在习惯行为的学习过程中奖励（积极的强化）和惩罚（消极的强化）的作用。如果一只关在笼子里的鸽子每次啄一下按

钮，就会得到一粒食物作为奖励的话，那么这只鸽子就会逐渐学习这种习惯。一旦这个习惯形成之后，只要食物持续不断地出现，那么鸽子就会继续啄按钮。但是，如果这种奖励停止的话，鸽子就会意识到之前的习惯不再起作用，因此，已经形成的习惯就不复存在。

行为主义心理学当中有一个有趣的领域，即用抵抗消亡的程度来显示习惯的长度与最初导致习惯形成的具体的强化程序类型之间的联系。问题是哪些类型的强化程序会导致产生最强的习惯。最简单的强化类型是对每一个行动进行奖励——每啄一下按钮奖励一粒食物。另外一种部分是强化，即仅在固定的时间区间给予奖励，例如每 60 秒一次，或者是固定的行为区间，每啄 10 下奖励一粒食物。研究显示了部分强化的强化程序会导致形成习惯更慢（习惯需要更长的时间才能形成），但是其维持的时间会更长。

大多数与技术分析相关的事实表明，几乎所有部分强化的强化程序、随机强化，¹²¹ 使得一种习惯更有抵抗力，而不至于逐渐消失。在随机强化的条件下，生物体的行为与它对接受奖励的反应之间没有真正的关系。某种行为受到奖励是偶然的。在一项实验中，让一只鸽子啄一个彩色的按钮，并且对它进行了时间为 60 秒的随机强化。结果是，在结束奖励之后，啄按钮的这一行为又持续了 3.5 小时。换言之，这只可怜的鸽子做出的无意义的行为所用的时间，是其通过随机强化形成习惯¹²² 所用时间的 210 倍。

随机强化是那些相信虚幻的技术分析相关性的人们所能够接受的最佳奖励方式。有时，也许是在偶然的情况下，信号 / 形态是会产生预期预测的结果的，因而会加强人们对这种方法的信念。行为主义者可能会预测这种观点将会有极强的抵抗力拒绝这种观点的消亡。迷信和沉迷于赌博的人，则非常依赖于虚幻的相关性，而这种习惯却是无法改变的。¹²³

当我在美国证券交易所工作的时候，我遇见了很多表现出迷信行为的交易者。例如，有一个人每天都要系同一条领带，而另一个人则坚持从固定的门进入交易大厅。这些奇怪的行为都是从最初的一组偶然的行为开始的，而这次偶然的行为也正好符合了他们自身的预期结果，所以人们会通过随机强化的方法来使这一行为能够长期保持下去。

图形分析中的错误信仰

主观的技术分析师，不管他们使用什么样的具体方法，都会对视觉图形分析的有效性深信不疑。在金融市场数据中，价格走势图表被视为发现可以利用的指标的最佳方法。技术分析的前辈们运用图形发现某种法则的基本原理，而如今，他们仍然保留着这些最主要的分析工具。在这一部分中，我们将解释为什么依赖图形分析的坚定信心是错误的。

当然，这并不是说图形在技术分析当中毫无作用。它们是形成对市场行为进行可以检验的假说这一理论的灵感来源。然而，在对这些推测进行有力地、客观地检验之后，它们所留下来的也只是假设，而不是令人信服的知识。

对指标进行非正式寻找的可行性

人们需要，也有能力在自身经历中寻找秩序和意义。“因为人们具有普遍的适应性，所以经过演化的事物已经深深地植根于我们的心中。我们可以适应有序的现象，而那些随机出现的现象却无法适应。发现形态并确认各种形态之间联系的这种倾向是导致知识的发现与进步的主要原因。但是这其中的问题是值得我们注意的，即人们的这种倾向性是非常强烈的，而且有时是无意识的，所以，我们有时会在根本就不存在相关性的事物之间寻求一致性”。¹²⁴

直觉，这种非科学的思维是指既没有对表面的假象提出质疑，也没有对其给出明确解释的事物全盘接受的一种倾向。不幸的是，那些乍看起来非常明显的事物并不一定都是有效的。对于古人来说，太阳围绕着地球运转是非常明显的事儿。那么对于技术分析的前辈们来说，是金融市场价格决定了图形形态与运行趋势。正如古代的天文观察者所看到的虚构的星座轮廓那样。例如，狮子座、猎户座都是以天空中随意排列的星星为基础虚构出来的。同理，还有那些在股票和商品市场中根据曲折变化的价格走势图表而描绘出头肩形态、双重或三重顶以及其他形态的人也属于这一范畴。这些形态形成了主观图形分析的专业术语。

我们的神经系统对图形都感到紧张万分，不论这些图形是真是假。约翰·埃尔德博士（Dr. John Elder）¹²⁵ 是著名的预测模型与数据挖掘领域的专家，他认为这是“浮云中的兔子”效应。然而，现在有非常好的理由质疑图形分析的有效性，甚至对基于这种理论的知识的有效性同样可以提出质疑。我个人认为实际出现在我们眼前的趋势与图形形态通常都是由我们内心的贪婪所产生的幻象，尽管这种说法比较随意，但我还是对发现之中的规律抱有强烈的希望。

前段所述的内容并不意味着金融市场真的缺乏趋势与图形形态。实际上，确实有强有力的实证证明图形形态具有预测力。¹²⁶ 也就是说，对图形形态的视觉观察并不是发现或者证明图形形态有效性的适当手段。研究显示，¹²⁷ 即使是专业的图表分析专家也不能清楚地对由随机过程产生的有效的市场图形与虚构的市场图形形态进行区分。一种分析方法如果不能从虚构的图形中分辨出真实的图形形态，那么它就不能做出令人信服的真实的分析。同样地，一个珠宝商如果不能告诉他的客户真钻石与杂货店出售的钻石之间的区别，那么他就没有办法评估一条钻石项链的价值。

财务数据中的虚幻趋势与图形形态

统计学家亨利·罗伯茨（Henry Roberts）曾经说过，技术分析落入虚幻的形态和趋势的陷阱主要有两种可能的原因。“首先，绘制股票价格的一般方法只画出了连续的（价格）水平线，而没有画出价格变化的水平线，那么这条水平线就不会显示出一条人工绘制的形态或者趋势。其次，机会主义行为本身所产生的形态，会引起对趋势带有欺骗性的解读。”¹²⁸

罗伯茨认为，技术分析对于在随机游走中出现的规律性很强的相同的图形形态是非常重要的。¹²⁹ 根据定义，随机游走缺少真实可靠的趋势形态，或者是任何可以利用的结构。然而，罗伯茨的随机游走图形显示了包括头肩顶、头肩底、上升三角形、下降三角形、三重顶、三重底以及趋势通道等在内的图形形态。

你可以通过掷硬币的方式创建一个随机游走图形。假设以 100 美元为初始点，如果掷硬币的结果是头像，则增加 1 美元，反之，如果是背面

的话，就减少1美元。换言之，硬币的每一面都代表着价格的变动，用图形来表示就是价格变动的水平线。每一次对掷硬币的结果的模拟价格与之前初始价格的基础上增加随机变动价格的累积代数总和是相等的。如果你以初始价格300美元开始重复上述过程，你会对出现在你眼前的图形形态大吃一惊。我们发现有些图形的形态与金融资产的实际图形形态非常相似。很明显，通过掷硬币所得出的图形形态与趋势是不能预测其未来走势变化的。

由此，我们提出了一个十分重要的问题：如果不是图形分析师所绘制的图形能够在随机游走的过程中发现规律，并且知道这些图形根本不可能具有预测力，我们又是如何认为同样的图形形态有预测力呢？难道仅仅是因为这些图形是确实出现在实际的图形当中吗？

或许我们不应该相信罗伯茨的眼睛。他是个统计学家，而不是技术分析专家。如果一个专业的图形分析师能够像一位珠宝商分辨出真假钻石那样，指出出现在随机游走中的图形形态是假的，是无效的话，那么罗伯茨的随机游走实验就无法证明任何东西，除非他本来就缺少图表分析方面的专业技能。当然，专业的图形分析师会声称他们可以告诉你出现在随机游走中的图形与实际的图形之间的区别。毕竟，如果在实际市场图形中形成的是供给与需求、或者是市场情绪波动的结果，就像技术分析所言的那样，他们是完全有可能从随机游走图形¹³⁰出现的偶然性曲折走势中分辨出随机游走图形与实际走势图形之间的区别的。

遗憾的是，专业的图形分析师一直以来都没有能够说明真实的图形与虚构的图形之间的差别。阿迪蒂（Arditti）¹³¹为此进行了测试。他发给图表分析师们每人一份实际图形与随机游走图形相互掺杂的图表。结果显示，参加实验的图形分析师们不能清楚地告诉他这些图形的真伪，而他们分辨这些图形的方法只有猜。这一结果又在随机进行的一项非正式研究中得到验证。主持这项研究¹³²的是沃顿商学院¹³³的教授杰里米·西格尔（Jeremy Siegel）。他向那些自认为是图形分析专家的华尔街顶级股票经纪人展示了一张制图（见图2-18，其中4幅是真的，4幅是模拟的）。这一次，专家们依然不能区分真实的图形与模拟的图形。随后，我又对商学院刚好完成技

术分析课程的研究生进行了类似的实验。他们的准确程度与我猜想的一致。我用 X 表示图中真实的图形。

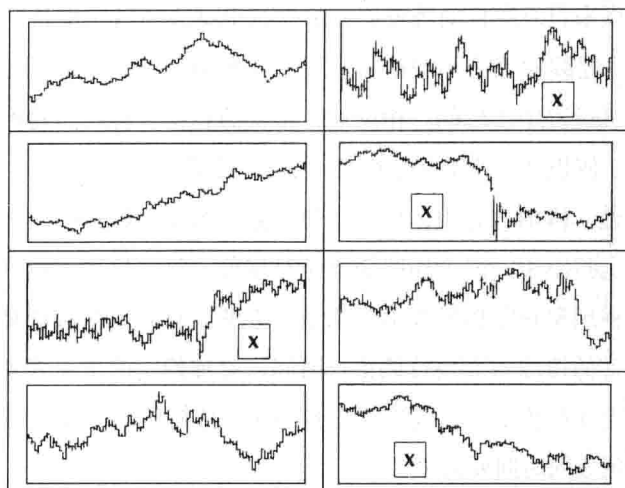


图 2-18 实际与随机产生的股票图形

资料来源：Siegel, Jeremy J., *Stocks for the Long Run*, 3rd Edition, Copyright 2002, 1998, 1994, McGraw-Hill; the material is reproduced with the permission of The McGraw-Hill Companies.

从以上的实验结果中我们得出了什么结论呢？首先，他们没有说明金融市场是随机的，并且因此缺乏有效的趋势和图形形态。但是他们确实对主观 / 视觉上的图形形态的分析是获得知识与预测的手段的有效性这一观点提出了质疑。图形解读是一种有效的方法，它最起码也应该可以从模拟的图形中分辨出真实存在的图形。

体育运动中的虚幻趋势的实证

不仅仅是金融市场里的观察者会受到顺序的错觉的困扰。许多体育爱好者和运动员也会相信成绩趋势的说法。即一段时间内持续的优异 / 糟糕的表现。棒球运动员和球迷并不认为击球手的连续击打不中或者连续的本垒打是可以实现的。棒球球迷、运动员、主教练以及体育评论员所说的所谓的“手感热得发烫”，是指击球手超过准确击打平均值的一种趋势。

人们广泛地持有这些观点，因为在竞技体育领域中，这种趋势表现得

最为明显。然而，在对选手的表现¹³⁴进行缜密的统计学分析之后，证明了这种手感火热的现象根本就是幻觉。球迷们一直被他们的直觉所蒙蔽。这种错误的观点来自对随机的现象应该怎样做出反应的错误概念。¹³⁵常识告诉我们击打（连续的三击不中被罚出局，或者出现连续的成功击打等）是不可能在随机过程中出现的。相反地，错误的直觉告诉我们随机过程可能会显示类似于使积极/消极的结果进行交互的东西。

总的来说，所谓的真相，就是认知错觉导致运动趋势的出现。我们所看到的连续击打不中，或者连续的本垒打现象，只不过是一系列相似的结果而已，这种现象在随机过程中是相当正常的。有效的趋势就是指比在随机数据中正常的击球水准要持续更长时间，更加稳定的击球表现。也就是说，击球状态（连续击中，或连续击打不中）所持续的时间比在随机过程中出现这种现象的时间要长。

吉洛维奇指出：“尽管与出于偶然性的打出好球的选手相比，普通球员并没有表现出持续的击球状态，但是这并不意味着球员的表现是由偶然性来决定的。某一次击打命中与否，是由非偶然的因素决定的，而最主要的因素是取决于进攻或防守球员的技术能力。然而，一个球员最近的击球命中率这一因素是无法对结果进行影响和预测的。”¹³⁶

那些在实际上只是随机的数据，是如何使得体育爱好者、运动员以及图形分析师落入虚幻的趋势和形态的陷阱当中的呢？这就需要我们对直观判断的本质做深入的研究了。

直觉判断与启发的作用

“简单来说，思维过程包括两种基本类型，无意识的思维过程和受到控制的思维过程。”¹³⁷“直觉属于无意识的思维过程。直觉是我们不需要对事物进行观察或者推理的过程，¹³⁸直接获取知识的能力。”正如普林斯顿大学的心理学家丹尼尔·卡尼曼所说：直觉是人们对知识与生俱来的、快速的认知能力。相反地，“经过深思熟虑的，或者受到限制和约束的思维过程具有推理、挑剔以及善于分析的特点”。¹³⁹受到控制的思维的原型科学推

理：直觉的原型来自通过嗅觉就能发现潜在病症的才华横溢的医学诊断专家。主观的技术分析师主要依赖于直觉，因此这种做法对于他们自身是极为不利的。

人类的智慧是有局限性的。我们仅仅对那些通过感官获得的零碎的信息给予太多的关注。此外，人们的大脑并不适合于在不确定的环境下，按照逻辑和可能正确的方式对复杂的事物进行评估并做出决策。

为此，我们必须将其进行简化。我们通过判断启发法（指一种快速，但存在瑕疵的推理方法）来降低评估和做出决策的复杂性。这种方法把我们的注意力限制在全部信息中的某一部分上，并且利用相对简单的信息处理和推理的方式，以最快的时间得出结论。我们所做的这些都是在无意识或者并不费力的状态下完成的。我们做出直觉的判断，就像识别某人的面容一样，是无意识的，也是毫不费力的，我们的大脑非常适合做出直觉判断。

我们从生活经历中获得判断启发法。启发法这一术语是指通过反复的实验来发现某种观点和想法。这也是获得判断启发法这种思考规则的途径。此外，由于我们在不知不觉中应用了这种方法，以至于最后我们不仅学习到了这种方法，并且接受了它们。究其原因，是因为这种方法已经被广泛地应用了。但是，当我们的判断受到失败的困扰时，这种方法就会暴露出缺点来。

启发式判断容易在连续的行为中出现错误。换句话说，判断出现了偏差——这种判断会以同样的方式重复发生错误。¹⁴⁰ 这种连贯性可能会对启发式判断的失败以及失败的原因进行预测——不论它与实际价值的背离是积极的还是消极的。

一般情况下，启发式判断在不确定的情况下容易产生错误。¹⁴¹ 金融市场存在高度的不确定性，因此在这个领域中，直觉的判断通常会产生错误。影响主观的技术分析师的具体偏见是对数据中心根本就不存在的趋势和形态存在偏好。换句话说，就是对只有在随机行为中才存在的感知秩序具有系统性的倾向。也许这也可以用来解释为什么图形分析师无法将真实的图形从随机数据中产生的虚构的模拟图形中分辨出来这一事实了。

在伯顿·马尔基尔（Burton Malkiel）教授的表述中，“随机性是人们非常难于接受的概念。当各种事件被归入不同的集群时，人们开始寻求解释和图形形态。他们拒绝接受经常出现在随机数据中的图形形态与来源于抛硬币所形成的图形形态相一致的这一事实。这种情况也同样出现在股票市场”。¹⁴² 马尔基尔认为市场是随机的这种观点是极端的。我虽然并不认同马尔基尔的观点，但是我相信人们会对随机数据中的秩序产生误解。正因为如此，对价格走势图中的视觉观察往往是不能让人相信的。

对随机数据中趋势和图形形态产生的错觉也许要归结于对具体的判断启发法应用的失败，我们将这种现象称为**典型性推理**。也就是说，对平时行之有效的典型性规则的错误应用，会使我们对本不存在的秩序的感知能力产生偏差。

启发式偏差和启发的有效性

概括地说，启发可以帮助我们在人类智慧受到限制时，对事物做出快速并且复杂的决策，但是这样会对所做的决策产生偏差。启发式偏差的概念可以简单地通过**启发的有效性**来解释。

我们依靠**启发的有效性**来评估事物未来发展的可能性。启发的有效性是基于合理的概念的，即我们越是容易想起某件事情，这件事情就有可能在将来的某个时间发生。我们把这种容易让人想起的事情称为**有效认知**。飞机失事就是与有效认知相关的事件。

启发的有效性具有一定的意义。能够对某类事情进行回忆的能力是与这些事在过去发生的频率相关的。一般情况下，在过去频繁发生的事情，在未来发生的概率就较大这种说法也是正确的。这个观点与概率论所阐述的“一个事件在未来出现的可能性与其历史的出现频率是相关的”这种观点是相一致的。¹⁴³ 同理，暴风雨在过去的发生频率比小行星撞击地球的发生频率要高，所以，暴风雨在未来出现的概率也更高。

伴随着启发的有效性这个概念出现的问题是，某些因素可以加强对事件的有效认知，而这些因素又与事件的历史发生频率没有任何关系，所以，它对事件未来的预测就显得不那么重要了。结果，由于受到这些不相干因

素的干扰，我们对可能性的判断有时会出现错误的夸大。当近因效应与生动性这两个因素与事件发生的可能性没有关系时，反而会增加对事件发生的可能性的有效认知。也就是说，最近¹⁴⁴正在讨论中的某类事件是如何发生的，它的逼真程度如何，以及我们是如何受这二者的影响将这些事件以最短的时间记入大脑的。假设飞机失事这类事件。一架飞机最近失事的事件对人们印象深刻。那么，刚好在一起众所周知的飞机失事事件之后的这段时间里，人们会倾向于高估未来空难发生的概率，并且因此对飞行产生过度恐惧。请注意，在判断中出现的偏见仅仅是过高地估计了众多可能性当中的一种。启发的有效性从来不会导致我们对某一事件发生的可能性做出低估。在这中间产生的错误通常是过度估计的可能性中的一种。

代表性启发：相似性推理

代表性启发，最初是由特沃斯基和卡尼曼¹⁴⁵发现的，它与主观的技术分析有着重要的联系。我们用这一规则进行直观的分类判断。换言之，我们通过此法来评估特定目标的分类。例如：我面前的一只狗，是属于贵妇犬这一特殊类别的。代表性启发式的基础前提是某一种类别中的成员就当显示出该种分类的主要属性。我们对同一类别中的对象应当拥有与本类别相近特点的看法是合理的。当我们通过具有代表性的判断来分辨一个事物是否属于某一种特殊的类别时，我们会下意识地认为物体间的相似程度，或者物体间属性特点的匹配，是这种类别的代表特征。这就是代表性启发式简称为启发式的原因。

由于人们对完全随机产生的（实际上是真实的数据）顺序和图形形态的认知能力存在偏差，所以我们在这一部分及以下的章节中，要对视觉的图形分析是有自身缺陷的观点进行讨论。我个人主张这种偏差是由于代表性启发式的不当应用所导致的。换言之，由于对代表性启发式的不当应用而产生的偏差是可以解释为什么主观的技术分析师会落入实际上是由随机过程产生的虚幻的图形顺序的陷阱这个问题的。

为了解释图表分析师为什么会犯这种错误，我们会通过由代表性的推理导致的启发式偏差的案例进行说明。

就像我们之前提到的那样,代表性的推理是以物体或者事件可以按照类别的固定特征进行分类的概念为前提的。而某一类型的固定特征可以理解为物体或者事件拥有该类别主要的特点。通常这种推理的方法很有效。然而,当某一事物的显著特征没有包含其所属分类的明显标志时,这种方法就会失效。换句话说,当最吸引我们注意的特征与对其所属类别进行判断的工作毫不相干时,¹⁴⁶代表性的推理就会不起作用。遗憾的是,直观的判断大多是以一个事物或者事件最为明显的特征为基础的,而不论这一类别的真实特征是否得到了真实的体现。

代表性的推理与其说是通过某一事物最为突出的属性与其所属类别最为突出的属性之间的匹配程度来决定这一事物的分类,不如说是在评估这种可能性的概率。两者间匹配的数量越多,那么这种事物属于某一类别的概率就会越高。“当代的心理学家假设我们对银行的出纳员、女权主义者、微型电脑以及各式各样的东西进行分类(等级)的概念是通过我们对所相信的属性的认知来体现的,而这种概念也同时对这些事物的本质特点进行界定。”¹⁴⁷我们通过图 2-19 表示。在很多的实例中,这种快速简单的分类规则给出了最为精确的结论。如果这种方法无效的话,人们是不会使用这样的规则的。

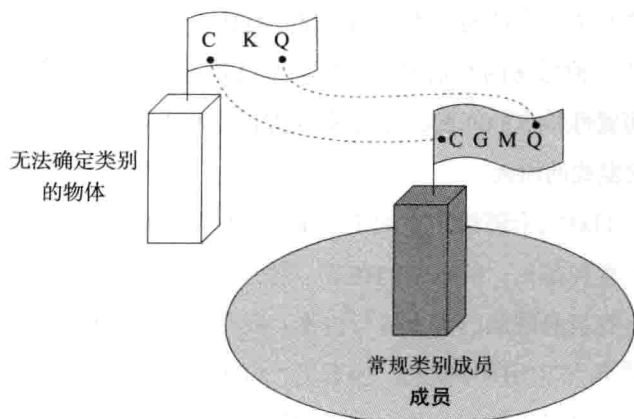


图 2-19 通过与属性匹配的数量来决定分类

我们现在来分析一个应用具有代表性的规则而导致对事物分类出现错误判断的案例。想象一下你已经上了一架商业航空公司的航班,而此时

你担心坐在你旁边的乘客（物体）可能是一名恐怖分子（类别）。因为这个人身上的明显的特征使你感到不安：这个人满脸的胡子，25～30岁的样子，是个典型的来自中东国家的人，头上包着头巾，正在读着《古兰经》，神情看上去十分紧张。这时，你会不由自主地认为这个男人肯定是属于恐怖分子这个类别的，因为这个男人的明显特征与发动2001年“9·11”恐怖袭击的恐怖分子的特征非常相似。你的推测就是以代表性启发为基础的。现在我们来分析另外一种情形，即如果这名乘客是来自中东地区的年轻男子，穿着蓝色的牛仔裤和T恤衫，但是T恤衫上面的图案却有伤风化，头上戴着印有纽约扬基队标志的棒球帽，而不是包着头巾，正在看一本《花花公子》杂志，Ipod音乐播放器中放着摇滚乐，嘴里还嚼着口香糖，还不时地吹着泡泡。面对这样的情况，你又会如何反应呢？当然，你不需要对此产生焦虑。人们对事物所属类别的直观判断是基于某种假设的，即个体的明显特征对于其类别的归属来说可以起到令人信服的指导作用。

基于某种代表性的判断会导致错误，因为它们没有把统计学/或者逻辑因素的情况考虑进来，而这种情况与对事物类别的准确判断紧密相连。对某一事件或者物体的类别进行评估的最可能的方法就是贝叶斯定理（Bayes' theorem），¹⁴⁸它是从概率论的原理中衍生出来的。如果将这一理论应用于怀疑飞机上的乘客是恐怖分子这个案例中，它会将世界上恐怖分子所占的比例这种因素考虑进去（例如关于恐怖分子等级分类的基本概率），即全世界拥有这种明显特征的（例如这种样本的基本概率）普通人的比例，以及恐怖分子中拥有这种特征的人数比例（即符合这些已知特征的条件，并被认为是恐怖分子的概率是多少）。当我们把这些数值代入表述贝叶斯定理的数学等式中时，通过假设一些人具有恐怖分子的特征，我们就会得到某人是恐怖分子的条件概率。

直觉是根据人们对某些有代表性的事物的理解而产生作用的，而不是以贝叶斯法则为基础的。对此，我们更倾向于犯一种名为**基本比例谬误**的错误。这种错误没有把描述某种情况具有哪些特征的基本比率的统计数字考虑进去。恐怖分子，是非常少的（也就是说，基础比率较低）。因此，在所有这些具有“9·11”事件中的恐怖分子相同特征的人当中，他们实际

上是恐怖分子的概率太小了。也许在 100 万人当中，才会有 1 个这样的人是恐怖分子。这个事实与我们判断邻座的乘客是危险人物的可能性事件高度相关。然而，我们仍然会忽视基础比率统计数字，因为直觉的判断往往聚焦在吸引公众注意的事件的特征上面，而不是像基础比率这样的相关事实上。

一个错误的代表性推理的例子是：市场分析师对伴随着 1987 年 10 月股市的暴跌而来的通货紧缩及经济大萧条的预测。他们做出如此预测的原因是，他们认为这次的股市暴跌与 1929 年股市崩盘后出现的通货紧缩及经济大萧条具有惊人的相似之处，即 1929 年和 1987 年股市都发生了暴跌。这种明显的相似之处使得分析师们把注意力从两者间的不同之处上转移了出来。也就是说，关于 1929 年通缩崩盘的预测信息与 1987 年的情况是完全不同的。对过去 100 年中股市崩盘案例的分析显示，价格的暴跌不能对通缩崩盘做出有效的预测。

另外一种与代表性推理联系在一起的错误，我们称为**结合谬误**。结合谬误没有把逻辑的基本规律考虑在内。它告诉我们的是一个真子集肯定比这个真子集所属的大型集合所包含的数量要少得多。例如：马是一个大型集合“动物”的真子集。而马又在“动物”这个广泛的内容中占很小的一部分，因为“动物”不仅包括马，还包括蓝知更鸟、土豚等。然而，当人们依赖有代表性的规则时，他们似乎忽视了逻辑的基本规律，并且承认了结合谬误。¹⁴⁹ 在一项由特沃斯基和卡尼曼¹⁵⁰ 进行的实验中，受试者会看到下面的这段描述：

琳达今年 31 岁，单身，为人坦率，也很聪明。她主修哲学。作为一个学生，她对有关歧视和社会公正的事情非常热衷，并且经常参加反核游行。那么下列的选项中哪一个是最可能发生的呢？①琳达是个银行出纳员，②琳达是个银行出纳员并且积极参加女权主义活动。

令人吃惊的是，有 85% 的受试者选择了第 2 个选项。这一结论虽然忽略了集合与真子集之间的逻辑关系。受试者如此肯定琳达具有明显的女权

主义特征的事实，实际上是犯了结合谬误的错误。很明显，女性银行出纳员和女性女权主义者这个集合，包含了在所有女性银行出纳员中存在女权主义者与非女权主义者。其他的研究也证实了这一发现，因此，特沃斯基和卡尼曼得出的结论是，人们的判断概率随着明显突出的细节数量的增加，而产生偏差。即使每一个额外描述事物的细节可以降低这个对象存在的概率。¹⁵¹但是当人们应用有代表性的推理时，额外的描述会增加与这个类别所匹配的特征的数量，因此，就会增加直观评估的概率。关于结合谬误，请参见图 2-20。

当结合谬误影响到陪审团的裁决时，就会产生灾难性的后果。例如，分析下面两种可能性当中的哪一种更有可能是真实的？¹⁵²



图 2-20 结合谬误

(1) 被告离开了犯罪现场。

(2) 被告因为害怕被指控谋杀，而离开了犯罪现场。

第 2 种假设包含了“离开现场”和“害怕指控”两个描述性的结合。因为第 2 种选择是第 1 种选择的子集，根据定义，它发生的概率很小。很明显，对于有些人来说，他们离开现场并不是害怕遭到指控，而是去取干洗的衣服。然则，第 2 种假设的额外细节增加了一层合理性，因为人们觉得这一细节与当事人的行为的匹配度大大增加了。

✓ 代表性启发式与错误趋势和图形中的形态：真实与虚幻

到目前为止，代表性启发式已经讨论了关于其在进行事物分类中的使用情况。实际上，启发式也包含在一个更加抽象的分类判别类型之中，而这种抽象的分类判别对于主观的技术分析来说是非常重要的：判断一个数据实证的样本的趋势和形态是否有继续分析的价值，或者说这个数据仅仅是随机的，从而没有继续分析的价值。

请记住，在判断一个数据集是否是以随机的形式出现时，并不是依靠自觉的意识和意图来决定的。它的出现就像我们辨认一个朋友的脸颊一样，是自觉而无意识的。所有的启发式判断，包括代表性推理，都是通过不自觉的、无意识的形式表现出来的。因此，当一个图表分析专家得出结论说，价格历史中包含启发式形态和趋势的样本（例如，数据是随机的）是值得进行分析的时候，他的表达就不是有关控制性思维的问题了。也就是说，图表分析专家无意识地提出了一个问题：“摆在面前的这份数据真的包含了可以开发利用的秩序和结构吗？”

尽管启发式判断在不知不觉中渗入了我们的思想，但是它们对我们的信念会产生严重的影响。最主要的问题在于启发式思维没有把数据的重要特征考虑进去，而这个数据又与决定它是否包含可以用继续分析的方法进行挖掘的启发式形态与趋势相关。

在前面的部分里，我要求你通过连续投掷硬币 300 次的方法创建一幅人工的价格历史走势图（随机游走图形）。如果你还没有做的话，我强烈要求你现在就做。因为这样做不仅会使这部分内容更加有意义，而且可以永久地改变你观察图表的方法。我们援引玛莎·斯图尔特（Martha Stewart）的话：“这可是一件好事儿啊。”

如果你是第一次接触随机过程产生的数据，你会非常惊讶地发现这份数据看起来与非随机数据竟是如此的相同。你也许也没有想过像头肩形态、双重顶形态等类型的图形形态会出现在大众面前。我们必须要考虑一个问题：为什么我们众所周知的由随机产生的排列顺序会表现出具有真实趋势和形态的特征呢？或者说，为什么通过纯粹的随机过程产生的数据看起来与通过法则创建的数据如此相似呢？那么随机产生的数据具备预测的能力吗？

你所看到的图形类似于图 2-21。

其本质是：看上去非

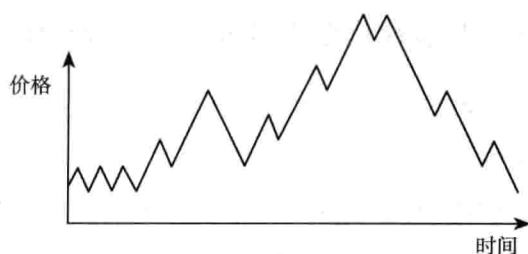


图 2-21 通过揭示排列顺序的随机过程所产生的数据

随机特征的图形并不一定就是非随机图形。按照我的观点，由于对代表性启发式的错误应用，人们认为在随机产生的数据中存在趋势和形态的感觉就会发生。进一步讲，启发式造成了一种错误的直觉概念，即应当如何看待随机数据。也就是说，启发式造成了错误的随机固定模式。因此，当我们发现与这错误的固定模式不相匹配的数据时，我们会不自觉且错误地做出数据是非随机性的结论。这一结论对错误的形态认识和相信形态中包含有预测力的信息的错误概念来说，无疑是一次飞跃了。

代表性推理的一种典型结论就是对结果应该与原因接近的预期。这种观点是可以理解的。地球（大的原因）是产生万有引力（大的结果）的原因，而鞋子里的石子只会带来极小的烦恼。然而，结果应该与原因更为接近的理论有时候是不成立的。例如，1918 年全球爆发的流感疫情，是一个大结果，而导致其发生的原因只是一种亚微观大小的动物。

那么以上所述的内容又与资产价格图形有什么关系呢？一个图形中的一组数据集合可以被视为过程的结果，即金融市场的活动过程。这个过程是原因，而数据是结果。我们的身体在 24 小时内的每一个小时的体温所构成的图形就是身体新陈代谢所导致的结果。

由于代表性思维是在自觉意识之下无意识产生的，所以当我们看到资产价格的历史走势图的时候，我们更加倾向于对产生图形的过程的本质做出下意识的结论。也许，如果股票的走势图像图 2-22 一样，技术分析的前辈们就不可能从中发现可以利用的结构，那么，技术分析的应用也就不会诞生了。数据被认为是一种接近于直观随机性的固有模式：它是偶然的、杂乱的，且不包括任何趋势的线索。

在代表性启发式的基础之上，我们假设如果一个过程（原因）是随机的，则它的结果（数据）应当也是随机的。作为这一假设的结果，任何与假设的随机性的固有模式相背离的

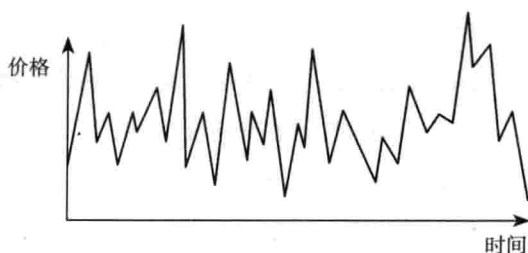


图 2-22 直观的随机数据的固有模式：偶然的、杂乱的、非趋势的

数据,因为显示了形态和趋势的特征,其结果也必定是非随机过程的产物。股票的走势图显示了由随机过程产生的轮廓。换句话说,这些特征符合了直观的排列顺序和可预测性的概念。请牢记金融市场在某种程度上来说也是非随机性的,因此这个市场也是有预测力的。然而,这并不是什么问题。相反地,有预测性的非随机的结构能否通过视觉的观察和直观的估计被发现出来呢?

像前面所讲的那样,代表性推理的缺陷是:在对明显的特征过度关注的同时,忽视了重要的统计学和逻辑特点。对此我们将在第4章进行解释。从数据分析的角度看,数据样本的一个最重要的特点,就是观察资料中所包含的数据样本数量。要想得到数据样本的有效结论,必须把观察资料的数量考虑进去。此外,大量的观测数据的统计学特点需要清楚地表现产生数据的过程所具有的本质特征(随机的或者非随机的)。相反地,小型样本数据经常表现出虚幻的图形形态和趋势,这是因为小型样本数据的统计学特点对数据产生过程不具有广泛的代表性。

我个人认为视觉图形分析是发现趋势和形态的有效方法这种观点是基于错误的思维链,比如下面的例子。

(1) 基于代表性启发式,认为随机过程产生的数据应该被看作随机的。这个观点是错误的想法,即一个没有任何形态、组织形态或者趋势提示的无形的、杂乱无章的触发器。

(2) 这种固有的形式,被误认为要体现所有随机的数据样本,而与包括样本在内的观察资料的数量无关。

(3) 对小型样本数据进行分析,发现它与直观的随机性的固有模式并不匹配。

(4) 数据因此被认为是非随机过程的产物。

(5) 非随机性过程具有预测性,因此在数据中寻找形态和趋势并做出预测,就显得十分合理了。

以上的观点如图2-23所示。

这种思考链的主要问题是没将样本规模(上述清单中的第2条)考虑进来。假设数据的样本可以清晰地反映数据生成过程的本质,而不考虑

组成样本的观察数据的数量，就会违反大数法则这种正规的数据分析法。代表性的推理没有遵守大数法则。换句话说，它承认犯了忽视小型样本数据或者样本规模的错误。这样做的后果就是，代表性的推理会不知不觉地成为傻瓜随机俱乐部¹⁵³的会员。

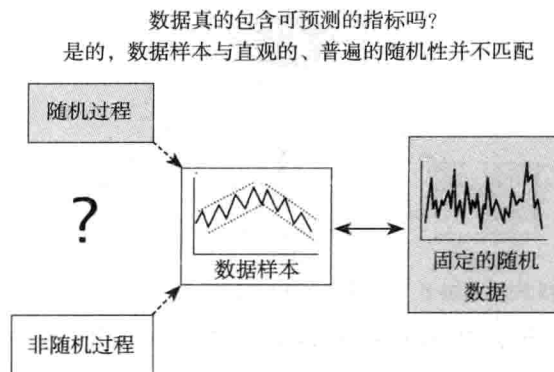


图 2-23 假设数据可以具有预测力，是因为它与直观、固定的随机性不相匹配

大数法则和忽视样本规模的错误

大数法则告诉我们样本的规模是最为重要的，因为它决定了我们基于数据样本所做出的推论的信任程度。大数定律，即大数法则告诉我们只有包括大量观察数据的样本才值得信赖，并且能够清晰地反映产生样本的过程的特征。或者，换种说法，只有大样本可以预示来自一个群体的样本的特征。所以，掷硬币的次数越多，就越接近在样本中出现正面的比例，这反映了在公平掷硬币的情况下会出现的潜在事实，即出现硬币正面的概率是 0.5。从本质上来讲，说明我们对从大样本，而不是小样本中学习知识的信心大大地增强了。在第 4 章中，我们将深度讨论大数法则与统计学的联系。目前，我们讨论的最主要的问题就是大数法则与小型数据样本中出现的虚幻的形态和趋势之间的联系。

小型样本容易引起误导，因为它可以证明与产生样本的过程的真实特征不同的特征。这个观点也解释了为什么在一个假设的股票图形中的一小部分（实际上是由随机过程产生的）可以出现逼真的形态和趋势。当然，这种图形是不可能出现真正的形态和走势的。因此，对于数据样本是否

出自随机，或者非随机的可靠结论，是不可能从小型样本当中获得的（见图 2-24 和图 2-25）。当图形的一部分（小样本）看上去是非随机的，但是当这些小样本结合在一起的时候，那么在整个图形中就会产生非随机性的幻觉。

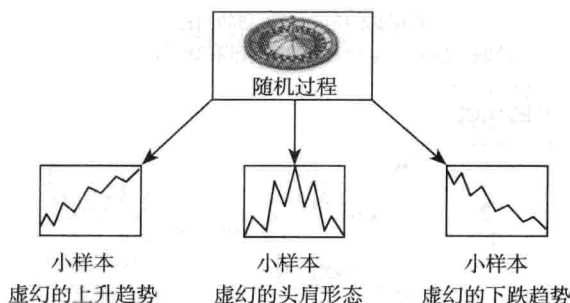


图 2-24 小样本中出现的虚幻的趋势和形态——由随机过程产生的数据

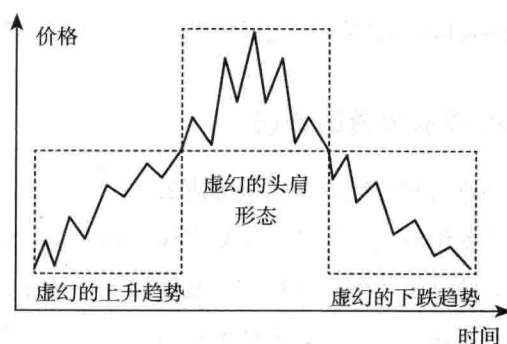


图 2-25 如果一个图形的一小部分看起来很有秩序，那么整个图形则呈现出非随机性的特征

掷硬币的试验显示了小样本是如何产生错误的印象的。随机过程最明显的统计学特点是会出现两个可能相等的结果。对于出现正面的概率是 0.50 的预期就是基于此的。当掷硬币的次数很多时，这种直觉是正确的。在一个掷硬币 10 万次的样本中，出现正面的比例非常接近 0.50。如果有人再次重复这个实验，出现正面的比例也是非常接近 0.50 的。实际上，这个比例的区间是 $0.497 \sim 0.503$ 。¹⁵⁴

直觉很少关注样本的规模或者大数法则。我们仍然期望在所有样本中硬币正面的出现概率是 0.50。尽管对任何样本中出现硬币正面的预期比例

通常是 0.5，即使是在掷两次硬币的情况下，在所给出的小样本中实际出现硬币正面的比例可能会与 0.50 这个数值严重的背离。例如，在掷 10 次硬币中，很容易出现 3 次正面（0.30）或者 7 次正面（0.70）。在掷 10 万次硬币的过程中，出现 3 万次（0.30）正面或者 7 万次（0.70）正面是根本不可能的。

其他关于随机过程的错误直觉是其结果存在可替代性（上升/下跌，正面/背面）。当人们想象掷硬币会出现哪种结果的排列顺序时，他们所想出的排列顺序要比出现在随机序列中的正面/背面的排列顺序多得多。¹⁵⁵在现实中，随机过程显示的可替代性结果和同一结果的连续性（上升、上升、上升）要比人们认为出现的结果少得多。因此，当一组数据不能与交替序列的错误预期相匹配的话，就会出现虚幻的趋势。例如，从直观的角度来看，掷硬币可能会出现正、反、正、正、反序列的可能性要比出现正、正、正、正、正、正序列的可能性看上去更加合理。然而，这两种序列在实际中都是有可能出现的，¹⁵⁶ 所以我们就会产生前一种序列是随机的，而后一种序列则是提示了趋势的错误印象。

因为随机序列实际上很少出现交替，而且也不像人们直观概念中想象的那样排列，所以，实际上是随机的序列反而以非随机的形式出现了。例如，棒球运动员的击打成功率是 50%，这与掷硬币实验中出现正面的概率相同。因此，在一场比赛中连续 20 次的击打过程中可能会出现连续 4 次、5 次，甚至是 6 次的绝佳击打机会。换句话说，球迷所说的“手感很热”的现象，¹⁵⁷ 实际上不过是随机过程序列的一种代表性的排列。

随机过程应当具有很强的交替性，而没有明显特征的错误概念是由两种错误的认识所导致的：①**聚类错觉**；②**赌徒的谬论**。文中，术语**聚类**是指时间或者空间，甚至是相似的事件所组成的群或者簇。在随机时间序列中出现的相似结果的簇，导致了趋势会继续下去的错觉。赌徒的谬论同样是一种错误的预期，他们认为一个由类似的结果组成的簇本来就不应当发生，但是当事实确实发生的时候，他们就会产生另外一种错误的预期，即转运已经离自己不远了。赌徒的谬论可以通过下面的例子说明，即赌徒认为在连续出现硬币的正面之后，就应该会出现硬币的背面了。

以上所讲述的是一个观察者对于自己是否落入聚类错觉或赌徒谬论的陷阱这一判断过程的观点。但实际上他并没有发现这个随机过程。如果这个人继续持有过程是非随机的，并且应当显示趋势这种观点的话，那么他就很有可能落入聚类错觉的陷阱。或早或晚，当排列出现的时候，之前的观点就会得到确认，即观察者已经落入了聚类错觉的陷阱之中了。

反过来，如果观察者认为过程是随机的，那么他就落入了赌徒的谬论这个陷阱了。例如，当连续4次出现硬币的正面时，那么接下来出现的就应该是背面了。尽管在随机过程中这种现象十分普遍，但是背面的出现会增强这种错误的观念。因此，当真实的随机过程出现时，认为这个过程是非随机的，并且相信趋势（聚类）将会出现的观察者和认为这个过程是随机的，并且相信这个趋势会出现反转的观察者都会认为自己的观点是正确的。其实，当这个过程是随机的，两种观察者的预期都是错误的。随机游走没有记忆功能，之前的结果对于后续的事件是没有任何影响的。

聚类错觉出现在一定的空间背景之下。在随机数据中形成的虚幻的概念是关于随机性的错误预期的另外一种简单表现形式。这种错误的预期认为随机数据不能显示任何结构或者形式。任何与此预期相违背的情况都可以被错误地解释为是非随机过程在工作的信号。我们来看一个案例，在随机数据中形成的虚幻的空间形态会促进错误观念的形成。在第二次世界大战中，德国使用 V-1 型和 V-2 型导弹轰炸伦敦的军事行动中，¹⁵⁸ 当时的报纸绘制了将要遭受导弹攻击的地点的位置图。市民们看到这些以后马上意识到这是密集轰炸的信号。但是在轰炸之后，事实证明所有的预测都是错误的，因为这种攻击形态是德国人为了避免攻击伦敦的某些地区而努力工作的结果，而这一结果又反过来导致了错误的因果关系推理。伦敦人开始相信这些地区之所以会免受 V-2 型导弹的攻击，是因为这些地区住着德国间谍。然而，对导弹攻击地点的正规分析显示了攻击形态与随机形态是完全一致的。

在加利福尼亚，促进错误的因果推理的虚幻空间聚类的案例是所谓的“癌症集群歇斯底里症”。如果某种特殊类型的（由石棉引起的肺癌）癌症的确诊病例超过这个社区的平均值，人们就会非常警惕。结果显示这种集

群的发生只是偶然性的。那么对单纯的统计学来说，其中并不明显的内容是：在对加利福尼亚进行 5 000 份人口普查，以及 80 种可能引起癌症的环境因素的调查中，由于偶然因素导致癌症发生的病例要高出平均值很多。在某些情况下，这样的集群确定与环境因素有关，但是也不总是像，或者完全像直觉暗示的那样。

现在让我们回到技术分析中来看一看双重顶、头肩底等有序排列的形态。这些形态的出现与通过随机过程产生的数据的直觉概念并不一致。这些数据被错误地认为是完全偶然的，并且是没有形态的。所以我们草率地下了结论，即形态必须是非随机过程附带产生的结果，而这个非随机过程不仅可以进行视觉分析，而且形态本身也要对预测有用处。正如罗伯茨所说的那样，随机游走可以很容易地追溯到被主观的技术分析所重视的形态当中去。

最后，人们很难说清楚一组数据在什么时候是随机的，或者不是随机的，因为人类的大脑倾向于接受有秩序的事物，并且善于编造故事来解释这些秩序的存在原因。我们也不难理解技术分析的前辈们在价格走势图中发现形态和趋势，并且为了这些形态、趋势的存在提供理论支持。我们要运用比视觉分析和直觉更为严谨的方法来发现在波动的金融市场中是存在可以被利用的法则。

解决虚幻的知识的方法：科学的方法

我们因为受到了欺骗而接受了错误的知识，本章对此进行了多方面的分析。到目前为止，解决这一问题的最好方法就是运用科学的方法，即第 3 章的主题。

第3章

Chapter 3

科学的方法与技术分析

技术分析的主要问题是错误的认知。根据传统的经验，许多技术分析可以说是一个在信仰和奇闻轶事的基础上形成的教条与谎言的结合体。它是通过技术分析的从业者使用非正式的、直观的方法来发现具有预测力的形态而得出的结果。作为一种准则，技术分析经常饱受诟病，因为它的应用实际上是根本没有准则的。

采用严谨的科学方法可以解决这个问题。科学的方法并不是可以自动发现知识的固定程序，而是一组“旨在描述并且解释被观察、被推断的现象的方法。它的目的是建立起一个可以检验的知识体系，并以此对事物进行排除或者确认。换句话说，科学的方法是以测试要求为目的，并对信息进行分析的具体方式”。¹本章总结了科学的方法的逻辑和哲学基础，并且讨论技术分析的从业者运用科学方法的含义。

最重要的知识：获得新知的方法

“在西方世界带给世界的所有知识中，最有价值的莫过于科学的方法了，它是一套获取新知的程序。这些知识是从公元 1550 年到 1700 年由众多的欧洲思想家创造的。”² 与非正式的方法相比，科学的方法将事实从谎言当中区分出来的能力是非常卓越的。在过去的 400 年中，我们对于自然界的理解与控制取得了长足的进步，这充分证明了科学方法的力量。

科学的方法的严谨性使我们免受智力缺陷所带来的困扰，这种缺陷会对我们用非正式的方法从经验中学习到的知识造成损害。尽管对于日常生活中许多明显的事实来说，获得非正式的知识会让我们认为它是真实合理

的，但是有时候我们认为明显无误的事情是错误的。这一点从古代的人认为太阳围绕着地球运转的观点中便可初见端倪。为此，这就需要科学来提示它的错误。当某些现象处于复杂或者高度随机的状态时，非正式的言论与直觉的推论更加容易遭受失败。金融市场行为就将这两个特点完全显示了出来。

历史上，技术分析还从来没有以科学的方法实践过。但是如今，这种情况得到了改变。在学术界和商界中，随着金融工程师这一新兴职业的出现，科学的方法就一直被应用着。实际上，一些最为成功的对冲基金也正在使用着一种被称为科学的技术分析策略。

不出所料，许多传统的技术分析从业者一直抵制这种改变。因为习惯的、倾向于既得利益的思维方式是很难被遗弃的。当民间医学进入到现代医学中时，当炼金术促进了化学的发展时，当占星术促进了天文学的发展时，出现反对的声音就不足为奇了。因此，传统观念的实践者与科学的实践者之间的矛盾就会显现。然而，如果历史具有指导意义的话，那么传统的技术分析最终只能是越来越被边缘化。而占星师、炼金术士以及巫医仍然会不断实践，但是他们的所作所为将不再被认为是重要的了。

希腊科学的遗产：喜忧参半的结果

希腊人是第一个致力于科学研究的民族，尽管后人们对希腊人的科学遗产的评价喜忧参半。从积极的角度来说，是亚里士多德发明了逻辑学。他发明的标准的推理过程是今天的科学发展的方法的核心内容。

从消极的角度来看，他创立的物质运动理论是错误的。按照现代科学的标准，亚里士多德的理论既不是从以前的言论之中归纳总结的，也不是通过对新鲜的言论进行检验而得出的，而是全部从形而上学的原理中推导出来的。当某种言论与理论相悖时，亚里士多德及其理论倾向于歪曲事实，而不是改变或者抛弃这种理论。

亚里士多德最终意识到演绎逻辑并不足以让我们认识世界。他意识到一种基于归纳逻辑的实证方法，即观察之后的归纳。这种方法的发明是亚

里士多德对科学的重要贡献。然而，他却没有把这项发明付诸实践。亚里士多德在雅典创办的高等学府——吕克昂（Lyceum）学园里，与他的学生们仔细地观察了大范围的自然现象。不幸的是，根据亚里士多德学派的教义，他们在没有充足的证据³的基础上，对事实做出的推论是存在偏差的。当事实与人们支持的理论相悖时，亚里士多德会将事实进行歪曲，以此来保证其理论的权威。

最终，亚里士多德教派留下的遗产在事后被证明是对科学发展道路的阻碍。但是由于亚里士多德本人的权威太高了，以至于在接下来 2 000 年的时间里，他的这些错误的理论被当作不可怀疑的教理由人们传承下来。这些都阻碍了科学知识的发展，或者说至少是那些具有科学特征的知识的发展。⁴同理，像道、沙巴克（Schabacker）、艾略特和江恩这些技术分析的前辈的学说，也是在没有受到质疑和检验的情况下传承了下来。不论是在科学中，还是在技术分析中，都不存在教义的说法。

科学革命的起源

“科学是 17 世纪最重要的发现，或者说是发明。那个时代的人都很博学——他们知道如何用科学的方法对自然现象进行测量、解释和处理，这一点可以说是具有革命性的发现。自从 17 世纪以来，科学取得了长足的发展与进步，并且发现了很多真理，人们从中受益匪浅。这些在 17 世纪都是不可想象的。但是，直到今天，科学都没有找到新的方法去发现大自然的真理。正是由于这个原因，17 世纪有可能会成为人类历史当中最重要的一个世纪了。”⁵

公元 1500 年左右发生在西欧的科学革命是在知识停滞不前的气氛中开始的。当时，所有的知识全部都是基于独裁者的声明，而不是来自观察的事实。罗马教廷的教义以及希腊科学的信条被当作千真万确的真理。在陆地上，物体受到亚里士多德物理学的控制，而在天空中，由希腊的天文学家托勒密（Ptolemy）发现的定律被认为是正宗的法则，并且受到了教廷的支持。托勒密的天动学说观点认为，地球是宇宙的中心，太阳、星星和行

星都是围绕着地球的轨道运行。对于非专业的观察者来说，这个事实看上去确实与教廷的正统理论相一致。

当人们开始注意到与这些真理相矛盾的事实出现的时候，人们表现得异常兴奋。例如，炮兵观察到由弩炮发射出去的炮弹与从加农炮中发射的炮弹的飞行轨道与亚里士多德的运动理论所描述的不一致（见图 3-1）。通过反复的观察，炮弹是以圆弧路径的方式运动的——这一结果与希腊人所做的轨道预测理论大相径庭，因为按照希腊人的说法，物体只要一离开投射的装置，就会直接掉在地上。以上的两种说法要么说明原有的理论是错误的，要么说明士兵们是受到了自身感官的欺骗。检验一种理论的有效性，就是要把这个理论的预测力与对其后续的观察相比较，这对于现代的科学的方法来说是非常重要的。因此，我们也可以把它看作科学的方法发展的第一步。

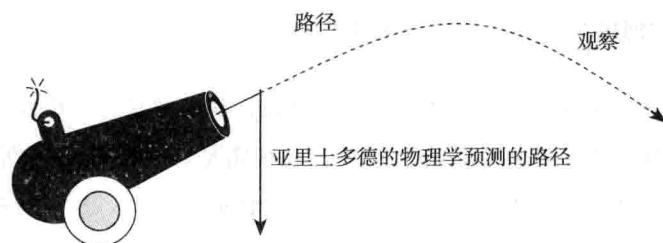


图 3-1 预测与观察的比较

同样地，在当时对于天空的观察也是与教会所持有的地心说理论相冲突的。大量的精确的天文测量仪显示，行星并不是按照教会所说的那样进行运动的。在最初的一段时间内，这些困惑的天文学家能够对教会所持有的理论进行修补，并使其能够与观察到的行星轨迹相一致。但是当不一致的观察出现时，原有的理论又增加了新的假设即每个行星都是围绕着一个较小的圆周运动。我们称为“本轮”（epicycles），并且用这个理论来解释棘手的天文学观察。例如，“本轮”有时可以用来解释为什么有时会出现反向运行（反向转动）。随着时间的推移，以及对“本轮”这个概念的不断修正，把从根本上就是错误的地心说理论转变成为了一个本轮之中包含着更小的本轮这样一个复杂的问题。

伽利略（Galileo Galilei, 1564—1642）发明的望远镜发现了科学史上具有里程碑的事件——木星有4颗卫星。这一发现与教廷持有的天上的物体都是围绕着地球运转的正统学说严重矛盾。伽利略的发明呼吁人们摒弃那些根深蒂固的信条，推翻教会对现实的看法。然而，宗教当局却不敢正视这些事实。尽管伽利略起初获得了教会的批准，可以将他的研究成果出版，但是后来教会收回了这个决定，并且判处伽利略有罪。作为惩罚，伽利略被迫放弃了他的天文学事业，并整日生活在教会的监视之下。⁶

最终，教会关于天体的理论被更为简单的哥白尼模型所取代。哥白尼认为太阳是宇宙的中心，并且试图用更少的假设对天文学的观察进行解释。如今，我们知道太阳并不是宇宙的中心，但是在哥白尼的那个时代，他的模型确实增加了人们对天体知识的认知。

对客观现实的信心与客观观察

自从科学对现实的本质做出了基本假设，我们发现有一种独立于、并且存在于观察之外的客观现实，这种客观现实是观察者感官印象的来源。这个信仰是科学观察的基础。它之所以具有信仰的特点，是因为现在没有任何实验或者某种逻辑能够证明客观现实的存在。例如，我自己并不能够证明我对红色交通信号灯的期望是由于存在于我自身之外的红灯所引起的，而不是在我大脑中完全浮现出来的某些东西所引起的。

也许有人会问，为什么我们会对那些作为常识来说是如此明显的事实大惊小怪呢？每一个人都会想象到一个存在于外面的世界。然而，在科学的眼中，这类问题是十分重要的事情，因为科学可以指出科学的认知与直觉的认知两者之间的差异。科学并不简单地接受那些看上去合情合理的事物。科学的历史不断地向人们揭示了显而易见的事物并不一定都是真正正确的道理。科学假设了一个独立的客观现实的真实性，不是因为现实能够不证自明，而是因为它与形成大多数科学组织的原则——简单性原则（principle of simplicity）相一致。

简单性原则认为越是简单的就越是最好的。因此，当所有的理论都符

合被观察的事实的时候，我们在判断众多理论中哪条才是正确时，简单性原则强烈要求我们接受最简单的理论。这种理论就是提出用最少、最简单的假设的简单性原则。简单性原则，也称为奥卡姆剃刀（Okaum's Razor），是指去除一个理论中多余的复杂性。哥白尼的宇宙模型要比教会的模型更为可取，因为哥白尼用一颗行星只有一条运行轨道的说法代替了教会对本轮的复杂叙述与解释。同理，如果市场历史中的某一段时间既可以解释为随机游走，又可以解释为波浪中又包含着波浪的艾略特波浪形态、黄金分割线、斐波纳契数列以及许多其他复杂事物时，在这种情况下，随机游走理论是最容易解释的。也就是说，除非艾略特波浪原理比随机游走理论能够做出更为精准的预测，否则，随机游走理论仍将是最佳的选择。

假设一个外部的现实的存在可以使复杂的事情大幅简化。其他方式则需要更多复杂的假设来解释为什么两个人观察同样的现象，会得出相同的印象。要想解释为什么我看到交通信号灯变红的时候，其他汽车也会停下来的现象其实很简单，即我假设在这里确实存在交通信号灯。然而，如果没有那个假设的话，那么在解释为什么当我停车时，其他的汽车也停下来的现象，就需要更复杂的假设了。

作为信仰客观现实的结果，科学对客观与主观的认知有着本质的不同。即因为我上班要迟到了，对红色交通信号灯的客观反应与我上班迟到的主观上的焦急状态，在这种情况下对红灯的反应是完全不一样的。客观的观察可以共享，并且被其他的观察者所确认。正是由于这个原因，客观的观察适合于建立可以共享，并且被其他人确认的知识。主观的思维、解释、情感以及其自身的缺陷，都不足以证明主观的技术分析是合理的知识。

科学的知识本质

阿尔伯特·爱因斯坦曾经说过：“在漫长的生活中，我学会了一件事：所有违背现实的科学、测量都是落后的、幼稚的——但是，它又是我们拥有的最为宝贵的财富。”⁷ 科学的知识不同于通过常识获得的知识、信仰、权威以及直觉。因为这些差异的存在可以证明科学具有更高的可靠性。这

也是这部分内容的重点所在。

科学的知识是客观的

科学通过把自己完全限制在客观的世界范围内，力求客观性的最大化。尽管我们明白获得完全客观的知识是可遇不可求的事情。客观的知识可以消除主观评估的顾虑，这种主观评估是一个人与生俱来独有的特征，并且容易被自己所接受。内心思想与情感状态是不能与其他人准确分享的，即使是艺术家、诗人、作家，还有作曲家的意图也是如此。威廉·巴勒斯（William Burroughs）在他的著作《裸体午餐》（*Naked Lunch*）中试图传递他个人吸食海洛因的经验。然而，我却从来没有感受到他的经历，我对那些信息的理解与威廉·巴勒斯所描述的经历简直是大相径庭，或者说与其他阅读过此书的人的感受是完全不一样的。因此，科学的知识就是指能够与人们共同分享并且尽最大的可能得到人们的认可的观点。这样一来，科学的知识就会促进独立的观察者们对某些事物的看法得到最大限度的一致。

科学的知识是以实证和观察为基础的。它是有别于从一套公式中所衍生出来，并且与其保持一致性的数学和逻辑推理的，而这些公式是不涉及外部世界和观察确认的。例如，勾股定理认为，在任何一个直角三角形中，两条直角边 a 和 b 的平方之和一定等于斜边 c 的平方，即

$$c^2 = a^2 + b^2$$

然而，这个真理并不是通过对上千个直角三角形的研究推导以及大量的观察所产生的，反而是从一套被人们普遍接受的数学推理中衍生出来的。

科学的知识是量化的

希腊的数学家毕达哥拉斯（Pythagoras，公元前 569 年—公元前约 475 年）提出了万物皆数的概念。正是这个论点认为世界和人的大脑在本质上是数学的。把重要的科学进行量化再怎么强调也不为过。“人类一直以来都有对事物进行测量的能力，也就是把它们转化为或者说是简化成为数字，这样的话，我们就可以在对事物的理解与控制上取得巨大的进步。然而，人类却一直也没有发现这样一条途径来衡量事物，更不用说在对心理

学、经济学和文学评论方面取得成功了，因为对这些方面的解释只是片面的，从而也就无法获得科学的地位。”⁸

如果要以一种理性的、严谨的手段进行分析的话，进行观察的资料就必须简化成为数字的状态。量化使得应用统计分析这一强有力的工具成为可能。“大多数科学家曾经说过，如果你不能够把你正在研究的东西用数学的方法描述出来，那么你就不是在做科学研究。”⁹ 量化是保证知识的客观性，以及最大化其与人们共享并被检验的能力最好的方法。

科学的目的：解释与预测

科学就是以发现预测新数据的法则，以及解释以前言论的理论为目的的。预测法则，经常是指科学的定律，是对反复出现的过程的描述，例如，“事件 A 可以预测事件 B”，但是它并不能够解释这种现象发生的原因。

解释性的理论要比预测性的法则更进一步说明事件 B 会发生在事件 A 之后的原因，而不是简简单单地告诉我们既有的事实。第 7 章我们要描述一些在行为金融领域非常超前的理论，这些理论试图解释金融市场行为有时候是非随机的原因。同时，这些理论也给予了技术分析以希望，并且对某些技术分析方法可以良好地进行做出了解释。

科学的定律和理论相对于它们自身的普遍性来说是有所不同的。被人们最为看重的事物往往是最为普遍的——也就是说，它们预测 / 或者解释大量的自然现象。对所有的市场及时间跨度都有效的技术分析法则与仅仅能够对每小时的铜期货数据做出预测的技术分析法则来说，它显然要具有更高的科学地位。

定律相对于自身的预测力来说也是有区别的。那些能够描述事物最为一致关系的预测被视为是最有价值的。对于其他的事物也同样适用于此。一个成功率达到 70% 的技术分析法则肯定要比成功率只有 52% 的技术分析法则有价值。

科学的定律重要的类型是函数关系。它把一系列的观察资料以等式的形式做出总结。这个等式描述了我们是如何对一个变量进行预测的（因变量），我们用字母 Y 来表示，而一个或者是多个被称为预测因子的变量的函

数，我们通常用字母 X 来表示。这个公式就是

$$Y=f(X_i)$$

通常预测因子的数值是已知的，但是因变量的数值是未知的。在许多应用中，因变量 Y 是指未来的结果，而 X 变量指的是当前已知的数值。一旦函数关系被推导出来，对于因变量来说，其预测值就可以通过代入已知的预测数值而得出。

函数关系可以通过两种方法进行推导，即从解释性的理论或者是用函数拟合（回归分析）的方法对历史数据进行评估而得出。目前，由于技术分析的理论正在逐步形成，所以技术分析主要受到后者的制约（详见第7章）。

逻辑在科学中的作用

科学的知识理应得到尊重，其中的部分原因要归结为它的部分结论是以逻辑为基础的。正是依靠于逻辑和实证来证明科学的知识的结论，科学才能够避免两种由于非正式推理造成的常见谬误：对权威的呼吁以及对传统的呼吁。对权威的呼吁是指对传说中的有学问的人的说法给予证明，而对传统的呼吁是指以长期以来固定的做事方法为榜样。

对于大多数流行的技术分析来说，谬误的不利影响主要归结于对于传统或权威的依赖，并非是正规的逻辑和客观实证。在很多的实例当中，当前的所谓权威仅仅是对其前任说法的引用而已，反过来说，对引用早期的专家的权威进行追溯，这些最原始的资料来源只不过是直觉产生的知识罢了。所以，权威以及传统所产生的种种谬误会因此而得到不断的加强。

逻辑的第一个作用：连贯性

亚里士多德（公元前384—公元前322年）被认为是形式逻辑的发明者，形式逻辑是从几何学中衍生出来的。埃及人在对线条和角度，以及对面积的精确测量方面已经有超过2000年的历史了，但是“希腊人把这种基础的概念进行延伸，并且把它们转化成为从数学定义（原理）中衍化而

成的令人瞩目的确凿的结论体系”。¹⁰

形式逻辑是与进行正确的推理法有关的数学的分支，而这种推理法是用来阐述和对辩论的有效性进行评估的。与非正式的推理相反，如果形式逻辑的法则可以进行后续的推导，那么得出正确的结果就会得到保证。

形式逻辑最基本的原理是连贯性的法则。它体现出了两种规律，即排中率（Law of the Excluded Middle）和矛盾率（Law of Noncontradiction）。“排中率是指在肯定、否定之间必须选择其一、不能两者全部否定。排中率不可能出现第三种可能。或者换句话说，中间的选项是被排除在外的。”¹¹ 对一个问题的表述非错即对，而不能出现两者都不可的情况。然而，排中率仅仅出现在二进制的应用当中，并且在不是真正的非错即对的情况下，是很容易被错误应用的。

“与排中率紧密相连的是矛盾率。它是指在同一思维过程中，对同一对象不能同时做出两个矛盾的判断，即不能既肯定它，又否定它。”¹² 在同一时间对同一事物的结论要求思想前后一贯，不能自相矛盾。

在接下来的这部分内容里，这些逻辑的规律将会对科学产生极大的影响。尽管在对诸如技术分析法则的盈利性进行回溯测试中得到了一些观察实证，但是并不能就此从逻辑上证明技术分析法则具有预测力，而同样的实证可以从逻辑上对技术分析法则缺少预测力的判断进行反驳（否认）。通过矛盾率，技术分析法则确实包含预测力的说法就会得到间接的证明。这种证明实证规律的方法，我们称为间接证明法或者矛盾证明法，它们是科学的方法的逻辑基础。

论点与论据

逻辑推理主要有两种形式：演绎法和归纳法。我们将对这两种方法分别进行说明，并且了解它们在现代科学的逻辑框架内共同作用的情况，即假设演绎法。然而，在考虑这些推理形式之前，我们先按照顺序把相关的概念了解一下。

- 论点。论点是指作为一种观点对某种事物是对或者是错的声明。例如，“头肩形态相对于随机信号来说，更加具有预测力”这句话就

是一种论点。论点区别于其他形式的陈述之处在于，论点不包含有像感叹句、祈使句以及疑问句那样带有错误的真理的特征。因此，只有论点是**可以被确认或者被否定的**。

- **论据。**论据是指一组论点当中的其中之一是作为结论出现的，而这个结论又是以逻辑的方式从其他的论点中得出的，我们称为**假设（前提）**。所以，论据声称假设为建立真理及其结论提供了有力的实证。

演绎逻辑

我们之前曾经提到过，逻辑推理主要有两种形式：演绎法和归纳法。这部分所提到的是演绎法；下一部分我们再来讨论归纳法。

1. 直言三段论

演绎的论据是指其本身提出的假设可以为真理以及其结论提供具有结论性的、确凿的证据。演绎的论据的基本形式是直言三段论。它是由两个假设和一个结论组成的。之所以这样命名，是因为它的作用是处理各个范畴之间的关系的。它首先会提出关于一个范畴的一般性真理的假设，例如，“人终有一死”这句话，在结尾处则会提出对一个具体例子所做的结论性描述，例如，“苏格拉底是会死的”。

假设 1：人终有一死。

假设 2：苏格拉底是人。

假设 3：苏格拉底是会死的。

请注意论据的第 1 个假设在两种范畴之间建立了关系，即人和死亡：人终有一死。第 2 个假设则对第 1 个范畴其中的一个具体的个人做出了描述，即苏格拉底是人。这就说明了这个个人一定是第 2 个范畴当中的一部分的结论，即苏格拉底是会死的。

直言三段论的一般形式可以用以下假设来描述。

假设 1：范畴 A 里的所有成员都是范畴 B 当中的一部分。

假设 2: X 是范畴 A 里的一部分。

假设 3: X 是范畴 B 里的一部分。

演绎逻辑有一个非常引人注目的特点——确定性，即通过演绎法得出的结论是具有完全确定性的真理，但是，有两个或者说是只有两个条件达成的时候：论据的假设是**真实的并且有效的**。如果其中的任何一个条件缺失，则结论就是具有完全确定性的错误。因此，对有效论据的定义是以它所做出的假设以及得出的结论是真实的为前提条件的。或者换句话说，对于有效的论据来说，同时拥有真实的假设和错误的结论是不可能的。

总的来说，通过演绎法得出的结论要么是**正确的**，要么就是**错误的**。如果任何一个有效的形式或者真实的假设缺失的话，所得出的结论肯定是错误的。如果两者都是真实的存在，则结论肯定是正确的。这里面是不可能包含中间立场的。

有一点是需要我们重点关注的，就是**真理和有效性**是两个明显的特征。真理，它的反义词是谬误，而谬误只是依附在某一个独立特征之下。如果一个特征与事实相符，那么它就是正确的。因为假设与结论都是论点，它们两者都具有正确与错误的特点。“所有的猪都会飞”这个假设是错误的。而“苏格拉底是人类”这个假设是正确的，那是因为我们得出了“苏格拉底是会死的”结论。

有效性是依附于论据之下的特征。换句话说，有效性是指包含论据在内的论点之间的逻辑关系。在对链接结论的假设的逻辑推理的正确性做出有效性或者其他缺失性的描述中，有效性并没有提及论据的假设或者对论据的结论的事实真相。如果说假设是正确的，那么它所得出的结论也必定是正确的，这时论据才是有效的。

然而，只要论点之间的逻辑联系是合理的存在，那么即使论点中包含有错误的信息，论据也可以被认为是正确的。下面的直言三段论具有有效的形式，是因为它的结论是通过**对假设进行强行的逻辑推理**产生的，而且它的假设与结论是明显不正确的。

所有的人都是不死的。

苏格拉底是人。

因此，苏格拉底是不会死的。

因为有效性与事实没有任何关系，所以有效性是可以由论据进行证明的，或者是通过论据图表来证明，我们称为欧拉圈，这个欧拉圈并不做出任何实际的推理。在一个欧拉圈中，一组要素包含一个范畴，我们把这个范畴用一个圆圈表示。从图中我们可以看到跑车被视为一个范畴，它是汽车这个更大的范畴中的次级范畴——汽车是整体，代表跑车的圆圈存在于代表汽车的大圆圈之中（见图 3-2）。

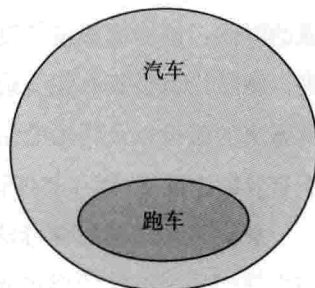


图 3-2 欧拉圈

图 3-3 更为清楚地解释了论据 1 是有效的，而论据 2 是无效的原因。论据 1 之所以是有效的，是因为它的结论，即“X 是 B”，是对假设的必要延续。然而，论据 2 就不具有有效性，因为它的结论“X 是 A”，不是对假设的逻辑推理。X 也许未必属于范畴 A。欧拉圈的描述要比论据自身更具有有效性、更有说服力，因为单独的论据还需要进一步的思考。

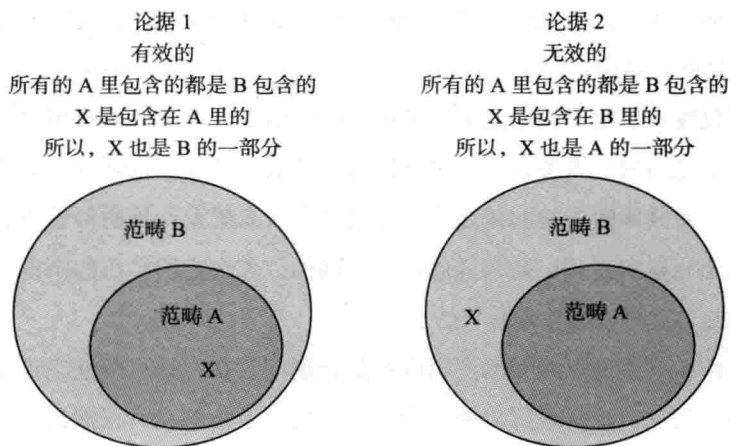


图 3-3 有效的和无效的直言三段论

论据 1

所有的 A 里包含的都是 B 包含的。

X 是包含在 A 里的。

所以, X 也是 B 的一部分。

论据 2

所有的 A 里包含的都是 B 包含的。

X 是包含在 B 里的。

所以, X 也是 A 的一部分。

有条件的三段论：科学论据的逻辑

对于科学的推理来说, 另一种形式的演绎论证就是**有条件的三段论**。它是确定新知识的逻辑基础。

与直言三段论一样, 有条件的三段论也包括三个论点: 两个假设和一个结论。以此命名的依据是它的第 1 个假设是有条件的论点。

有条件的论点是一种复合式的陈述, 它把两个简单的论点以“如果——然后”的形式连接起来。“如果”语句后面的论点作为**先行从句**, “那么”语句后面的论点作为**结果从句**。有条件的论点的基本形式是:

如果 (先行从句), 那么 (结果从句)。

例如:

如果它是一条狗, 那么它有四条腿。

在这个例子中, “它是一条狗”是先行从句, “它有四条腿”就是结果从句。第 2 个例子则是更接近于结果的表述, 即

如果技术分析法则具有预测力, 那么对其收益率的回溯测试值就会大于零。

我们这么做的最终目标就是要确定有条件的论点中的先行从句的真实性, 就像前面的第 1 个例子那样。尽管结果会逐渐清晰, 但是确定真实性的过程却是间接的!

有条件的三段论的第 2 个例子提出了一种论点, 即对第 1 种假设当中的先行从句和结果从句分别进行肯定或否定。我们以第 1 个样本为例, 基

于它的第2种假设可以是下面4种表述当中的任意一种。

它是一条狗：是对先行部分的真实性做出肯定。

它不是一条狗：是对先行部分的真实性做出否定。

它有四条腿：是对结果部分的真实性做出肯定。

它没有四条腿：是对结果部分的真实性做出否定。

有条件的三段论所做出的结论是对第1种假设的剩余的从句部分的真实性做出肯定或者否定。换句话说，结论所涉及的从句是在第2种假设中所没有提到的内容。例如，如果第2种假设所指的是第1种假设当中的先行部分，也就是“它是一条狗”这句，那么结论指的就是结果从句，“它有四条腿”。

下面是有条件的三段论的一个例子。

如果它是一条狗，那么它有四条腿。

它是一条狗（对先行部分的真实性做出肯定）。

因此，它有四条腿（对结果部分的真实性做出肯定）。

当有条件的三段论的组织构成是合理有效的时候，则结论从逻辑上来讲就是由它的两个前提所得出的。而且，如果有效的组织形态和它自身的前提是真实的，那么有条件的三段论的结论肯定也是正确的。

2. 有条件的三段论的有效形态

有条件的三段论有两种有效的形态：**确认先行项和否认结果**。在有关**确认先行项**的争论中，第2种假设对第1种假设的先行从句的真实性做出肯定。在对**否认结果**的争论中，第2种假设证明第1种假设的结果从句是不正确的。下面我们就将列举这两种有效的形态。它们都假设狗有四条腿。

确认先行项。

假设1：如果它是一条狗，那么它有四条腿。

假设2：它是一条狗。

有效的结果：所以，它有四条腿。

在这种有效的形态中，第2种假设通过说明“它是一条狗”，来确认先行项从句。这两种假设由此推论出“它有四条腿”这个结论。有条件的三段论的一般形式（先行项被确认）是：

假设1：如果A是正确的，那么B就是正确的。

假设2：A是正确的。

有效的结果：所以，B是正确的。

有条件的三段论的另一种形态是否认结果。在这种形态中，第2种假设确认第1种假设的结果从句是错误的。通过这一点，我们可以总结出第1种假设的先行从句也是错误的。请看下面的例子。

假设1：如果它是一条狗，那么它有四条腿。

假设2：它没有四条腿。

有效的结果：所以，它不是狗。

有条件的三段论的否认结果形态的一般形式是：

假设1：如果A是正确的，那么B就是正确的。

假设2：B不是正确的。

有效的结果：所以，A不是正确的。

我们也可以把它简化为：

如果是A，那么就是B。

不是B。

所以，不是A。

需要注意的是，论据的这种形式运用证据（生物没有四条腿）来证明先行项（它是一条狗）是错误的。在科学当中运用这种推理形式去证明假说的错误。假说只起到了先行项的作用。我们假设动物是条狗。如果假说是真实的，那么结果从句就会对将要被观察的东西进行预测。也就是说，如果这个动物确实是条狗，那么，当我们观察它的时候，就会看到它有四

条腿。换句话说,有条件的论据,如果假说(A)实际上是真实正确的,那么B就会在实验当中被观察。第2种假设表明在进行观察时,观察会对预测进行反驳,也就是说,B没有被观察。由此,我们可以有效推断假说A是错误的。

如果我们证明假说A是错误的,我们可以间接地证明其他假说是真实的,这一点我们会在适当的时候进行说明。其他的假说是一种新的知识,我们希望确定它是真实的,例如,有一种疫苗要比其他的(试验药物用的)无效对照剂更有效,或者说有些技术分析要比某些随机产生的信号的预测力更加有效。

3. 有条件的三段论的无效的形态

运用正规的逻辑的一个重要的理由就是解决在对有条件的三段论进行非正规的推理时,人们所遇到的困难,与我们在第2章讨论到的很多种偏差和错觉一样,心理学研究¹³显示人们在对有条件的三段论进行推理时,倾向于确认两种错误:确认结果的错误和否认先行项的错误。关于确认结果的错误,我们来看下面的例子:

假设1:如果它是一条狗,那么它有四条腿。

假设2:它有四条腿。

无效的结果:所以,它是一条狗。

而实际的情况是有四条腿的动物不一定就是有四条腿的狗。它只是有可能是狗。这里,我们看到证据与“它是一条狗”的事实相一致。然而,证据也会与其他的结果相一致。如上面的例子所提到的,这种动物有可能是奶牛、马,或者其他有四条腿的动物。确认结果的错误的一般形式如下:

假设1:如果A是正确的,那么B就是正确的。

假设2:B是正确的。

无效的结果:所以,A是正确的。

这种错误在不成熟的科学推理以及在很多已经得到肯定的技术分析中是最为常见的一种错误。下面的三段论,就是确认结果的错误的例子。

假设 1：如果技术分析法则 X 具有预测力，那么它就应当在回溯测试中产生利润。

假设 2：回溯测试是有盈利的。

无效的结果：所以，技术分析法则具有预测力。

假设 1 和假设 2 都有可能是正确的，但是结果却不一定是正确的。关于第 1 种假设“如果技术分析法则 X 具有预测力，那么它就应当在回溯测试中产生利润”，我们认为它是正确的。换句话说，能够产生盈利的回溯测试应该与技术分析法则具有预测力是一致的。然而，能够产生盈利的回溯测试可能会与技术分析法则仅仅是靠运气的假设相吻合。同样地，就像证据显示有四条腿的动物就是狗一样，它也和其他有四条腿的动物的假设一致。¹⁴ 确认结果的错误是实证不能够用于证明假说正确与否的原因。我们接下来要看到一位科学的哲学家卡尔·波普（Karl Popper），他声称科学的方法必须依靠于对结果的否定（错误），即我们之前曾经提到的推理的有效形式。

其他与有条件的三段论相关的错误就是否认先行项的错误。请看下面的例子：

假设 1：如果它是一条狗，那么它有四条腿。

假设 2：它不是一条狗。

无效的结果：所以，它没有四条腿。

事实上，不能把不是狗的动物排除在“有四条腿”这个范围之外。否认先行项的错误的一般形式如下：

假设 1：如果 A 是正确的，那么 B 就是正确的。

假设 2：A 是不正确的。

无效的结果：所以，B 是不正确的。

从逻辑上将其进行简化之后的形式是：

如果是 A，那么就是 B。

不是 A。

所以，不是 B。

当关注复杂的对象时，无效的论点很难被注意到。最好的方法就是在 A 和 B 当中插入普通的项目来显示。

图 3-4 和图 3-5 总结了有条件的三段论的预测结果。

有效的	无效的
确认先行项 如果是 A，那么是 B 是 A 所以，是 B	否认先行项 如果是 A，那么是 B 不是 A 所以，不是 B
否认结果 如果是 A，那么是 B 不是 B 所以，不是 A	确认结果 如果是 A，那么是 B 是 B 所以，是 A

图 3-4 有条件的三段论的一般形式

有效的	无效的
确认先行项 如果是一条狗，那么它有四条腿 它是一条狗 所以，它有四条腿	错误：否认先行项 如果是一条狗，那么它有四条腿 它不是一条狗 所以，它没有四条腿
否认结果 如果是一条狗，那么它有四条腿 不一定有四条腿 所以，它不是狗	错误：确认结果 如果是一条狗，那么它有四条腿 有四条腿 所以，它是狗

图 3-5 有条件的三段论的一般形式

正如我们所提到的，演绎推理的最大优点就是传递具有确定性的结果的真实性。然而，演绎逻辑有一个很大的弱点：它不能够揭示关于世界的新知。这是因为演绎的论点只能揭示存在于假设当中的真理。换句话说，演绎的推理仅能够检验存在于假设当中，但是还没有被证明的真理。

这并不意味着最小化或者忽视演绎，反而去说明其作用。在假设当中隐含的结果很有可能与明显表明结果存在很大的差异。伟大的数学发明

是最精准的——它所揭示的真理是隐含在数学体系的公理当中，在被证明之前还没有被人们所理解。这与人们广为宣传的费马最后的定理（Fermat's last theorem）相似。这一定理在 1665 年出版的一本书的空白页中被首次提及，但是直到 1994 年才被安德鲁·怀尔斯（Andrew Wiles）和理查德·泰勒（Richard Taylor）所证明。

归纳逻辑

归纳法是发现的逻辑。它的目的是通过超越包含在归纳的论据中的假设来揭示关于世界的新知识的。然而，这种新的知识具有很大的不确定性。通过归纳得到的结论本身就是不确定的。也就是说，它只是在某种程度或者概率上是正确的。所以，概率的概念是与归纳法紧密联系在一起的。

归纳法的行为方式与演绎法完全相反。我们看到演绎法从一个假设当中推导出一个普遍的事实，并将这个事实运用到众多没有限制的实例当中去，例如，从“人终有一死”到“苏格拉底是会死的”这个具体的事实。反过来，归纳性的推理跳出了以对数量有限的被观察实例为基础的假设，进而对数量无限但又没有被观察到的实例做出一般性结论。因此，这种推理形式通常是指归纳概括。它已经超越了通过观察而直接获得经历的范畴，并且会产生不确定性，即归纳法所得出的结论有可能是错误的。

归纳概括可以用下面的例子来表示。

假设：1 000 条狗里面的每一条狗都有四条腿。

一般结论：所有健康的狗有四条腿。

这个结论，就称为普遍概括，因为它肯定了这个等级中所有的狗都有四条腿。普遍概括的形式如下：

所有的 X 中所包含的东西都在 Y 的范畴内。

或者

X 中 100% 的东西都有 Y 的特征。

非普遍概括的形式如下:

X 中所包含的 P% 是属于 Y 的范畴。

或者

X 中所包含的 P% 有 Y 的特征。

或者

X 含有 P 比例的 Y 特征。

我们看到概率与概括的概念是分不开的。实际上,非普遍概括也称为概率归纳。例如,斗牛犬要比贵宾犬更具暴力性。这个统计学上的非普遍概括¹⁵既没有声称所有的斗牛犬都有暴力行为的倾向,也没有主张大多数的斗牛犬有这种倾向。然而,它说的是以类为单位,统计显示斗牛犬比贵宾犬的危险概率更高。

亚里士多德作为科学家的失败在某些方面要归结于他没有注意到非普遍概括所提供的有用的知识。他对演绎法的确定性的迷恋导致了他拒绝进行非普遍概括的研究。亚里士多德不能够理解自然界的重要规律在本质上是具有概率的。

列举归纳法

归纳的论据最普通的形式是以列举为基础的。它为包含在一组观察资料之内的假设列举实证证据,并且得出一般性结论,而这个结论是与所有在列举范围之外的类似的观察资料相联系的。

假设: 在过去 20 年的时间里,技术分析法则 X 有 1 000 次发出买入信号的实例,在其中的 700 次实例中,市场在接下来的 10 天内出现大涨。

结论: 在未来的时间里,当技术分析法则 X 发出买入信号时,市场在接下来的 10 天里出现大涨的概率是 0.7。

通过归纳法得出的任何结论在本质上都是不确定的,因为它扩大了观

察资料的范围,而这些资料还没有进行观察。然而,一些归纳论据要强于其他的论据,因此,它所得出的结论要更有确定性。归纳论据的优势与结论的确定性主要依靠对这种假设所提供的实证的数量和质量来决定。越多高质量的实证被引用,则用于未来观察的结论就越精确。

设想一下技术分析法则 X 的成功率是 70%,而这个值又是仅仅靠 10 个实例,而不是以 1 000 个实例为基础得出。在这种情况下,这种结论的不确定性最少增加了 10 多倍。¹⁶这意味着如果将来技术分析法则 X 发出的信号的准确性与其历史的成功率相比,出现很大的差异时候,我们就不会感到意外了。在科学领域,为支持结论所提出的实证是以统计学的方法进行估算的。这就允许我们对结论的不确定性做出大量的陈述。这个话题我们会在第 5 章和第 6 章中进行讨论。

归纳论据的优势还可以以实证的质量为基础。实证的质量完全取决于其自身,但是我们完全可以说有些观察性的方法确实能够得出更高质量的实证。科学中的黄金标准是核对实验,在核对实验中,只有一个被研究的因数是常量。技术分析不允许进行核对实验,但是我们可以利用不同的方法来进行观察。与技术分析有关的一件事是系统性错误或者说是系统性偏差。作为在第 6 章将要讨论的内容,如果不能进行谨慎研究的话,客观的技术分析研究倾向于发生一种特殊类型的系统性错误,我们称为**数据挖掘偏差**。

另外与归纳论据的优势相关的,是对有可能在未来碰到的观察资料的种类的典型性做出实证引用的程度。一个关于某种法则发出多头/中性(+1, 0)信号的推论显示,在市場上涨的时期内对正在处于市場下跌环境中的技术分析法则的绩效进行回溯测试的话,其准确性往往并不准确。即使这个技术分析法则没有预测力,它对多头或者中性头寸的制约使得它有可能产生盈利,这是因为该法则是在市場上涨的时期内进行的回溯测试。

归纳法最普通的错误是轻率地做出概括,即基于少量的或者质量较差的实证做出归纳。对一种技术分析法则的研究只引用数量很小的成功信号作为基础,这样对于具有预测力的法则做出的结论就有可能出现轻率的概括。

科学的哲学

科学的哲学试图理解科学发挥作用的原因及其工作方式。科学的哲学解释了如下事情：科学论据的本质，以及它与非科学和伪科学论据之间的区别；科学的知识产生的途径；科学是如何解释、预测，并运用技术利用自然的；确定科学的知识的有效性的手段；运用科学的方法的标准以及得出结论所用推理的类型。¹⁷

科学的方法是人类最伟大的发明之一，也是到目前为止最为有效的获取关于自然世界的客观知识的方法，这一点是毫无疑问的。“作为世界的组成部分，西方的工业化世界被视为科学革命的里程碑……”，¹⁸ 人类预测和控制自然世界的能力在过去的近 400 年中得到了很大的提升，科学的革命的开始，距智人（现代人的学名）出现在这个世界上已经有 15 万年了。

令人感到不可思议的是，科学的方法的发明以及它所取得的最初的成果，直到人们理解了科学的方法能够如此有效的工作之后才被认可。对科学的方法的领悟，历经了几个世纪的时间，随着从事研究的科学家和他们不断提出的批评，科学的哲学家不断对此进行钻研，最终理解了科学的哲学的工作原理以及它是如何工作的。

事实证明科学的方法所做的工作是远远不够的。科学的方法对于理解事物的本质是非常重要的。哲学家们所发现的让人烦恼的事情，都是明显矛盾的东西。一方面，牛顿的运动定律及万有引力原理和人类用不断发展的技术控制自然之间的矛盾。另一方面，科学的知识从本质上来讲是具有不确定性的，因为通过逻辑推理得出的结论是不确定的。那么这么多有缺陷的推理方法又是如何产生出这么多有用的知识的呢？

这一部分讨论的是科学的方法的发展，和对科学的哲学的工作原理以及它是如何工作的进行更深层次的理解的主要步骤。读者也许希望跳过历史的发展这一环，直接对科学的方法的关键因素进行总结。

培根的热情

哲学家们从来不需要别人的邀请，就会自觉地置身于所有的领域。这

就是他们的特点。他们告诉我们如何做事（行为准则），政府应该如何统治（政治哲学），什么是美（美学）以及与我们最紧密相连的事情，即构成有效的知识（认识论）的内容，还有我们应该怎样去获取它（科学的哲学）。

从某种角度来说，科学的革命是对亚里士多德学派科学的反抗。希腊人将物质世界视为不可靠的真理的来源。亚里士多德的导师——柏拉图认为，世界仅仅是对存在于“原型”（Forms）之中的真理和完美之事物进行错误的复制，在一种形而上学的非物质领域中，会出现的完美的狗、完美的树以及其他能够被发现的无法想象的事物。

当对这种现实的观点进行的反抗最终达成时，它所表现出来的是坚持不懈的努力所得出的结果。¹⁹ 经验主义作为思考的新学派，强烈地反对希腊人的传统形式。经验主义学派主张不仅自然世界是值得研究的，而且通过仔细的观察也是可以揭示事物的真理的。经验主义学派的先驱，也可以说是科学的哲学的第一位哲学家——弗朗西斯·培根（Francis Bacon，1561—1626）曾经说过：“本质，就是那种一目了然，不能够被没有偏见的思想所误解的东西。”²⁰ 在他的巨著《新工具》（*Novum Organum*）一书中，培根赞美了观察和归纳的力量。他认为科学将会成为一种完全客观的推理，通过没有偏差和偏见的观察，对知识进行肯定或者歪曲，然后从这些观察当中进行概括归纳。在大多数情况下，这种方法看起来是行之有效的。

然而，经验主义并不是培根及其追随者认为的最好的工具，而哲学家们则把解释这个观点当作自己的任务。首先，他们指出，经验主义是以观察为基础的，依赖于对通过观察所不能确认的事物进行重要的假设，他们假设本质在时间和空间里是统一的。这种假设对于证明如果一种科学的法则在此时被认为是有效的，那么其在其他的任何时点都是有效的这个观点是非常重要的。因为本质的统一这种假设是不能够通过观察来确认的，他只存在于信念之中。其次，科学通常处理那些用直接的观察所无法描述的现象和概念，例如原子的结构、重力的力量以及电场。尽管它们的影响是可观察的，但是它们的结构是不能够通过观察和归纳得出的。相对于可以进行观察的物质现实来说，上述的这些现象和概念作为具有预测和解释功能的人类的发明更容易为人们所理解。

然而，培根对于科学的方法的发展所做的贡献是十分重要的。它促进了尝试的思想，并且通过对不和谐的观察制定特殊的注释，为质疑提供了空间。

笛卡尔的质疑

如果说哲学家们样样精通，这个比喻对于勒内·笛卡尔（Rene Descartes, 1596—1650）来说是再适合不过了。被喻为现代哲学之父以及科学起源核心人物的笛卡尔，对于希腊的权威的知识和罗马教廷的教义采用了培根的怀疑主义。然而，笛卡尔对经验主义者关于观察的力量和归纳概括所持有的观点产生了质疑。他的传世名言“我思，故我在”表达了科学必须以怀疑任何事情作为开始的观点，而那些经历过质疑的人则不包括在内。此观点认为知识必须是通过演绎推理建立起来的，而我们所运用的推理应该没有受过有缺陷的人的五官造成的错误的观察干扰过的。

笛卡尔的反经验主义立场和其对事实真空理论的嗜好所导致的结果是，他的科学发现基本上是没有意义的。²¹ 然而，他对科学的贡献却是持久的。怀疑主义对科学的贡献是非常重要的。另外，笛卡尔发明的解析几何为牛顿发现微积分扫清了障碍，而微积分又为牛顿证明其著名的运动方程和重力提供了基础。

休谟对归纳法的批评

对归纳法进行质疑的还有来自苏格兰的经验主义者和哲学家大卫·休谟（David Hume, 1711—1776）。休谟的那部影响深远的巨著《人性论》（*Treatise on Human Nature*）出版于1739年，这本书主要解决的是认知论的中心问题：如何将知识从松散的状态中分辨出来，比如说某种观点碰巧成为了真实的存在。在休谟的著作出版之前，哲学家普遍认同将知识从松散的状态中分辨出来是与获取知识的方法的有效性有关系的。但是我们所要证实的是一种智慧的知识，而不是获取知识的方法的来源。²²

然而，哲学家们在哪种获取知识的方法是最好的这点上并没有达成共识。经验主义者辩称伴随着客观观察的归纳概括是获得智慧的最佳途径。

另一方面，理性主义者主张对那些不证自明的真理进行演绎推理才是正确的方法。

休谟对于以上两个思想派别的争论最终成为了他与自己的斗争。作为一名经验主义者，休谟鄙视理性主义者那种纯粹的演绎推理方法，因为这种方法与被观察的事实没有任何关系。而且，在经验主义的精神世界中，休谟认为调整人们对实证所持有的信任的比率是非常明智的，他还认为理论是要经过与之相匹配的观察来评估的。但是，当休谟对经验主义的逻辑基础进行攻击的时候，他陷入了自相矛盾的境地。休谟拒绝承认归纳概括法的有效性，并且对建立因果法则的科学能力不屑一顾。他对归纳法的批判被后人称为“休谟的问题”。

休谟对归纳法的攻击包括心理和逻辑两个方面。首先，他认为归纳法能够在事物之间建立相关的，或者因果关系的信念只不过是人类心理学的副产品。休谟肯定因果的关系这一概念仅仅是大脑产生的假象。A 引起了 B，或者与 B 相关的这个观点，是因为 A 通常在 B 之后出现，这只不过是习惯罢了。

从逻辑的角度来看，休谟声称归纳法是有缺陷的，这是因为它没有足够数量的观察实证，而且不论实证是否经过客观的选择，它都能够依靠有效的演绎推理得出结论。另外，他还认为，没有归纳法的法则会告诉我们，当我们有足够数量和质量的实证时，就可以证明那些超出有限观察范围的实例，而这些实例都是类似的无限大的、并且还没有进行观察的实例。一种归纳法本身就是基于早期的归纳法之上产生的有效的归纳法的结果，而这种循环会永久地继续下去。换句话说，试图证明归纳法完全依赖于无限的回归的观点——从逻辑上来讲就是不可能的谬误。

在休谟对归纳法的攻击之下，归纳法的支持者们节节败退，最后他们声称从狭义上来讲，归纳概括法只在概率的层面上是正确的。只要积累的实证能够证明 A 和 B 之间的关系，那么二者之间的关系就是准确可信的概率就大大增加了。然而，哲学家们很快就指出对归纳法进行概率层面的解释是有缺陷的。正如我们将在第 4 章中提到的那样，A 可以预测出 B 的这种概率等于 B 出现在 A 之后的次数除以 A 中的全部实例数量，不论 A

之后是否会出现 B。²³ 因为未来当中包含着无数的实例，这种概率的出现有可能会一直为零，它是与过去进行观察的实证数量无关的（任何数字除以无限大的数字的结果都是零）。

因此，休谟和他的追随者们创立了悖论。一方面他们表面上对归纳法的缺陷进行关注，另一方面是在进行令人震惊的科学发现的积累。如果科学是以这种带有缺陷的逻辑为基础的，那么它又是如何取得成功的呢？

威廉·维赫维尔：假说的作用

哲学家和科学家用了 200 年的时间观察科学的方法是如何工作的，以及理解它为什么会取得如此的成功。到 19 世纪中叶的时候，科学已经把归纳法和演绎法的逻辑有机地结合在一起。1840 年，威廉·维赫维尔（William Whewell, 1794—1866）出版了《归纳法科学的历史和哲学》（*The History and Philosophy of Inductive Sciences*）一书。

威廉·维赫维尔是第一位理解归纳法在假说的形成当中发挥重要作用的人。维赫维尔将假说称为快乐的猜测，他认为科学的发现开始于将所有大胆的归纳带入全新的假说当中，但是随之而来的却是演绎法。也就是说，在假说建立之后，就要从已经建立的假说当中进行演绎推理的预测。这种预测形成了一种有条件的陈述方式：

如果假说是真实的，那么对未来具体的观察就会按照预测如期的发生。

当进行观察的时候，假说也许会与预测的一致，从而确认假说的真实性，或者与预测的不相一致，从而与假说相矛盾。

什么是具体的假说？就是对怀疑的部分进行猜测，例如：

X 可以预测 Y。

或者

X 引起 Y。

这种猜测是在科学家早期的经历中形成的——注意 X 和 Y 组合的重复出现。一旦 XY 假说提出，那么可以检验的预测就会从中推理出来。这种预测形成了一种有条件的陈述方式，它包括两个从句：先行从句和结果从句。假说作为先行从句，预测作为结果从句。在 X 仅仅与 Y 相关联（预测）的情况下，那么条件的陈述形式为：

如果 X 可以预测 Y，那么将来 X 的实例后面就会出现 Y 的实例。

如果假说确认 X 引起了 Y，那么条件的陈述形式为：

如果 X 引起了 Y，那么如果 X 被排除，则 Y 就不会发生。

有条件的陈述中的结果从句里所包含的预测可以与新的观察资料相比较。这些观察资料带来的结果是未知的。这一点非常关键！当然，这也并不意味着观察资料就涉及将来的某些事件。它只是意味着当进行预测的时候，观察资料产生的结果是不可知的。在诸如地质学、考古学等历史科学中，与观察资料相关的事件是已经发生过的。然而，它的结果仍然是未知的。

如果对 X 未来的观察显示 Y 并没有出现在 X 之后，或者说由于 X 的排除阻碍了 Y 的出现，那么这个假说（先行）就被证明是错误的，我们可以通过有效的演绎推理证明结果的错误。

如果是 X，那么是 Y。

不是 Y。

有效的演绎推理：因此，不是 X。

然而，如果跟随 X 未来的实例出现的确实是 Y，或者说如果 X 的排除导致了 Y 没有出现，那么这个假说就是无法证明的了！记住，确认结果不是一种有效的演绎推理形式。

如果是 X，那么是 Y。

是 Y。

无效的演绎推理：因此，是 X。

对预测的确认仅仅是对假说的尝试性确认。此刻暂时存在的假说，接下来将会受到最为严格的检验。

维赫维尔对于休谟提出的归纳法的猜想是人类思维的一种习惯的说法表示赞同，但是休谟一向轻视他人的习惯使得维赫维尔的观点虽然有些许的神秘，但也算是富有成效的了。维赫维尔无法解释导致这种创造性思想的精神过程，即使他相信归纳概括法就是这种思想当中的一部分。他呼唤着一种能把假说变为不可传授的、富有创造性的、又非常关键的资质的能力。因为，如果没有假说的存在，一组观察资料所留下来的只不过是一组没有联系的事实，它们既不能进行预测，也不能进行解释说明。然而，直到有一个人实现了科学的突破。

维赫维尔对科学创意方面的关键描述使我回想起了若干年前我与百老汇最为成功的、最多产的音乐片作曲家查尔斯·施特劳斯（Charles Strauss）的一次交谈。在我 20 多岁的时候，我做事鲁莽，缺少社会经验，当时我问他是如何成为一位多产的作曲家的。而他的回答已经超乎了我的想象，查尔斯向我描述了这样一个略带宗教色彩的每日行为准则，即每天早晨 8 点到 11 点，下午 2 点到 5 点都要坐在钢琴的前面，不管你是否创作的灵感。他告诉我他把作曲视为他的工作。在这段时间里，他尝试着创作一些旋律——音乐的假说。他很谦虚地把自己的成功归结于 99% 的勤奋加上 1% 的创造天赋。然而，当我多年以后想起这段谈话的时候，心中又是另一种滋味。在几乎无限大的音符组合中，他可以创作出美妙的乐曲获得成功是可以称得上是一种天赋的，即拥有特殊的能力把那些潜在的对音乐的想象创作成为旋律，而这种能力其他人是没有的。这就是维赫维尔所说的快乐的猜测——这是一种不可传授的天赋，即把对于主旋律的领悟与一组与之没有任何关系的事物或者音符联系起来，而这一切是普通人无法达到的。这种能力只有少部分人拥有。

维赫维尔意识到一种假说的提出就如同搞一项发明，它与蒸汽机和电灯泡的发明有异曲同工之处，它们都代表了关于科学的思考²⁴方面最为重

要的发展阶段。归纳法本身并不能产生真理，但是它是获得真理最重要的第一步。这一点严重违背了之前对于科学的定义，即科学作为一种系统性的客观调查，伴随着它而来的是归纳概括法这一定义。维赫维尔认为科学家既是创造者，又是研究者。

卡尔·波普：证伪和演绎推理的科学回归

在《科学发现的逻辑》(*The Logic of Scientific Discovery*)²⁵和《猜想与反驳》(*Conjectures and Refutations*)²⁶两部具有里程碑意义的著作中，卡尔·波普(1902—1994)将维赫维尔对于科学的理解进行了延伸，并且通过对演绎法作用的解释说明重新定义了科学发现的逻辑。波普的中心观点是科学研究是不能够证明假说是真实的。相反地，科学的局限性在于无法分辨哪一种假说是错误的。为了解决这个问题，我们需要将实证的观察与随后进行的证伪中运用的有效的演绎推理形式相结合。

如果假说H是真实的，那么实证E在特定的条件下(例如，
技术分析法则的回溯测试)就会发生。

实证E在特定的条件下没有发生。

因此，假说H是错误的。

在此立场之上，波普对由某哲学流派提倡的一种时下非常流行的、被称为逻辑实证主义的观点发出了挑战。正如弗朗西斯·培根曾经反抗希腊传统的束缚那样，波普现在挑战的是逻辑实证主义的发源地——维也纳学派。逻辑实证论者相信通过观察是可以证明假说是正确的。波普对此提出了异议，他认为经过观察的实证仅仅可以用于证明假说是错误的。科学只是谎言的检测器，而不是真理的检测器。

波普证明了他的方法，即证伪主义，请看下面的例子。一组已知的观察资料可以通过大量的假说进行解释，也可以与这些假说保持一致。因此，其自身进行的观察并不能够帮助确定那些假说当中的哪一个是最接近于正确的。²⁷假设“我的一只鞋不见了”这句话。一种假说可以将其解释为我是一个经常乱放东西的人，而另一个假说则可以对“我的一只鞋不见了”

做出相一致的解释,即我的房子里闯入了一个只有一条腿的贼,因为他只需要一只鞋就足够了。实际上,有无限多的假说可以对丢失的鞋这件事进行前后一致的解释。

我们已经看到数据并不能够从逻辑上推测出假说的真实性。我们要尝试着去承认确认结果所导致的谬误。²⁸然而,这一点是波普的证伪方法的关键所在,数据可以通过否认结果有效地推导出一种错误的假说。换句话说,无效的实证可以用来揭示错误的解释。例如,另一只鞋的发现会证明一条腿的贼这个假说是错误的。那么合乎逻辑的论点应该这样来表示:

假设 1: 如果一个一条腿的贼是导致我的鞋丢失的原因的话,
那么我就不会在房间里发现我的鞋。

假设 2: 丢失的鞋找到了(否认结果)。

结果: 因此,关于一条腿的贼的假说是错误的。

波普的证伪方法与常识发生了激烈的碰撞,它在取代确定性实证方面存在偏差。正如在第 2 章提到的那样,直觉通常告诉我们在确定性实证能够被发现的时候,对假说的真实性做出检验。我们就这样处在确定性实证足以建立真理的错误印象当中。然而,我们正在以这种形式运用实证来证明确认结果的谬误。我们认为对“如果假说是真实的,那么确认性实证就应当被发现”的质疑是正确的,但是我们认为“确定性实证足以建立真理”的说法是不对的。换句话说,确认性实证是命题的真理必须具备的重要条件,但并不是充分的条件。²⁹波普的观点认为,重要条件的缺失足以建立假说的虚伪性,但是重要的实证的存在并不一定就能够产生真理。有四条腿是作为狗这种动物的重要条件,但是四条腿的存在并不能够证明这种动物就是狗。然而,观察资料所显示的没有四条腿的动物对于它是一条狗的这个论点来说,就是谬误。波普所持有的证伪的逻辑可以被视为反对影响非正常推理(见第 2 章)的确定性偏好的保障。

我们所举的最后一个实例说明了错误的实证的力量,以及确认性实证的弱点。那就是由著名的哲学家约翰·斯图尔特·穆勒(John Stuart Mill, 1806—1873)提出的黑天鹅事件。假设我们想要搞清楚一个命题的真实性:

所有的天鹅都是白色的。穆勒认为，波普也同意他的观点，即不论有多少只白色的天鹅被观察——也就是说不论有多少确认性实证的数量，这个命题的真实性是永远不会被证明的。一只黑天鹅也许会躲藏在下一个角落里。这就是让休谟感到心烦意乱的归纳法的局限性。然而，也许通过仅仅观察一只黑天鹅，人们就可以肯定地宣布这个命题是假的。下面所显示的有条件的三段论中，认为“所有的天鹅都是白色的”这个命题是以伴随着结论而来的证伪进行有效的演绎推理的形式为基础的。

假设 1：如果所有的天鹅是白色的说法是真实的，那么将来所观察到的天鹅也都是白色的。

假设 2：发现了一只不是白色的天鹅。

有效的结果：所有的天鹅是白色的命题是错误的。

根据波普的证伪法对一种假说进行检验的论点，其一般的形式为：

假设 1：如果假说是真实的，那么在未来的观察中，可以预测属性 X 的出现。

假设 2：观察显示属性 X 没有出现。

有效的结果：因此，假说是错误的。

以上的实例就是用于对统计假设进行检验的逻辑，我们将在第 4 章和第 5 章中进行讨论。

1. 科学的知识的本质：临时性与积累性

波普的证伪方法包含着某种暗示，即所有存在的科学的知识全部都是暂时性的。不论这种理论现在是否被广泛地接受，它都将成为实证所挑战的对象，而且这种理论被更加准确的理论所替代的概率是存在的。今天，爱因斯坦的相对论被认为是正确的。尽管它所预测的东西已经经受了无数的检验，也许明天就会通过一种新型的检验方法证明相对论是错误的，或者说不完整的。因此，科学永远是一个没有止境的推测、预测、检验、证伪以及进行新的推测的循环。科学的知识体系就是运用这种方法逐渐地向更加精确的客观现实的表述演变。

在大多数情况下，当陈旧的理论被替代的时候，并不是因为它们被证明是错误的，甚至是不完整的。当牛顿的运动定律被爱因斯坦的理论所替代的时候，牛顿的物理学在很大程度上依旧是正确的。在我们有限的日常生活范围内，物体是以匀速状态运行的（小于 90% 的光速），牛顿的这个定律仍然是正确的。然而，爱因斯坦的理论的正确性表现在更大的范围内，它不仅包括物体每天的运动，而且还包括光速。换句话说，爱因斯坦的理论，是建立在牛顿的理论基础之上的，它更加全面，并且包括了牛顿的理论。

以曾经成功的理论为基础建立起来的新理论（那些已经被确认的预测）和剔除错误的思想（那些被证明是错误的预测）共同构成了不断改进的知识体系。这也并不意味着任何被提出的新型知识学科就不一定是更好的。今天任何领域里的科学家比生活在上一代的科学家要知道得多的事实是无可争辩的。然而，黑帮说唱音乐和莫扎特的音乐之间孰优孰劣却是个永远争论不休的话题。

2. 从科学的限制到可检验的声明

波普的证伪方法的另外一层含义就是：科学必须把自己设定为可以检验的假说，即观察资料产生预测的论点并没有出现。也就是说，经过检验的假说能够继续存在还是被证明是虚假的，是通过对新的观察资料进行推理之后得出的结果是与预测的结果相一致，还是相矛盾得出的。运用新的观察资料对所做的预测进行比较，是培养不断改进的科学的知识最为重要的途径。正是出于这种原因，没有产生可进行检验的预测的论点必须排除在科学的讨论范围之外。

术语“预测”，在假设检验的情况下，是指澄清、说明的意思，因为技术分析从本质上来讲就是与预测有关的工作。当我们提到预测时，就会把它和对假说的检验联系在一起，当然，这也并不是说明对观察资料做出的预测在将来就一定是谎言。相反地，术语“预测”指的是对观察资料进行预测的最终结果是未知的这一事实。

在诸如地质学这种历史科学中，预测指的是对已经存在超过几十亿年的事物进行推测。例如，地质学中占主导地位的板块构造论学说，会对下

面的事情做出这样的预测：如果明天对某些特殊的地形进行科学研究的话，那么形成于几百万年前的地质结构就会被发现。而在金融领域，有效市场假说会做出如下的预测：如果对技术分析法则进行回溯测试，那么其经过风险调整后的利润，是不会超过经过风险调整的市场指数收益率的。

一旦从对假说的推测当中得到预测的话，那么产生新的观点所必需的行动就要执行了。这些行动包括对经过预测的地质结构的位置进行考察，或者对技术分析法则进行回溯测试。然后用观察的方法对预测进行大概的比较。对假说是正确的还是错误的观察、预测和做出决策三者间的一致程度的测量就是统计分析所包含的内容。

从科学的角度来看，最重要的一点就是假说能够对结果处于未知状态的观察资料进行预测。这个观点就使得对假说进行回溯测试成为了可能。因为观察资料总是要适应不同的解释，所以观察资料的结果是已知的这个说法就不能够满足假说的目的，我们所说的观点就会与既成事实相一致。因此，我们所做的努力就不可能，或者说不愿意做出可检验的预测，所以，这就为对观察资料的预测成为谎言的可能性打开了大门，这样一来，它们也就不再符合科学的标准了。

3. 界线划定的难题：科学与伪科学的区分

波普的证伪方法的重要结论就是：它解决了科学的哲学的核心问题，即对科学和伪科学的范围进行了定义。科学的范围局限于做出预测的命题（推测、假说、观点、理论等），而预测又是可以进行实证证伪的。波普所指的就是诸如可以进行检验的，或者是有意义的命题。不符合可以检验和有意义标准的命题，就是不可进行检验的和没有意义的。换句话说，这些命题没有说出任何有内容的东西，或者没有达成测试的预期。

不能被证明为假的命题似乎坚持着某些东西，但是实际上，它们却并不如此。不能被证明为假的命题之所以没有被质疑，是因为它们没有说明将要发生的事情的具体程度。事实上，它们并不包含大量的信息。因此，命题的可证伪性是与信息内容相联系的。可以被检验的命题的信息量是相当大的，因为它们可以做出明确的预测。即使可以被检验的命题被证明是错误的，但是它至少说明了某些实质的内容。不能够产生可以检验的预测

的命题在本质上是可以产生任何结果的。例如，天气预报说：“明天多云或者晴天，潮湿或者干燥，有风或者无风，寒冷或者炎热”，这种预测可以使得所有可能的结果发生。但它并不是谎言。我们对此能够得到的唯一结论，就是这种预测缺乏内容是非常明显的。

问题是伪科学家和伪预报员对他们进行预测的说话方式如此熟练，以至于可以掩盖他们在信息方面的缺乏和不能被证明为假的命题。假设占星术士预测了这样一件事情：你会遇到一位高大黝黑的陌生人，而你的生活也会因此得到改变。不论最后的结果如何，这句话都不可能反驳。你将会得到以下两种结果当中的一种：①耐心等待——你很快就会遇到这个陌生人；②你已经遇到了这个黝黑的陌生人，而你的生活也确实不同于以往，你自身也发生了明显的改变。这两种回答是不可能受到质疑的，因为两种预测是非常模糊的。它是不能表明你的生活是什么时候开始发生改变的，或者通过哪种明显的（真正的）方式实现的。“在下个星期三的晚上7点之前，你会见到一个穿着红色鞋子的人，走在第42大街并且吹着歌曲‘丝绸娃娃’的口哨”，像这种表达与预测相比相去甚远。“下个星期三的晚上7点”这个情形将会出现，你将有机会来评估这些预测的真伪。即使做出的预测最终被证明是错误的，至少这种可以被检验的事情可以让你决定将来去拜访占星术士是否还有意义。

出现在深夜电视购物节目里的推销员是表述没有意义的言论的最佳人选。“戴上我们的铜手镯，会提高你打高尔夫球赛的成绩。”这句话中“提高”是指什么意思？我们通过什么方式，在什么时候对这个预测进行检验？就像占星术士所做的预测那样，这种言论的模糊性使得它不可能推导出一个可以测试的（可以被检验的）预测。然而，我们很容易想象在既成事实发生之后的那些奇闻轶事是有可能得到确认的。“天哪！戴上铜手镯之后我确实得到了放松。过去，我在高尔夫球赛开始之前，总是感觉到紧张，现在，我的情绪稳定多了。上周我几乎还打出了一杆进洞的成绩，现在我也不像以前那样总是满腹抱怨了。我的妻子也说我穿上高尔夫球赛的比赛服显得英俊多了，而且我觉得我也不再掉头发了。”尽管没有意义的言论会导致各种奇闻轶事，并且使人们对此表示信任，但是这样做会阻止客观的

可以被检验的证据的出现。说这种话的人是不会被人抓到把柄的，他们的言论反而会得到那些被确认的奇闻轶事的报道的支持。

反过来，“戴上铜手镯会使你击球的距离提升 25 码^①”这句话，是含有丰富的信息，并且是有意义的，因为它能够产生可以进行测试的预测。你在没有戴铜手镯的情况下击打了 100 次球，然后计算它们的平均距离。第二天，在戴上铜手镯的情况下，你又重复了昨天相同的运作，并再次计算它们的平均距离。现在，你把相同的动作连续做 10 天。如果戴着铜手镯时得出的平均距离比平时少于 25 码的话，那么反驳这种言论的证据就掌握在我们的手中了。

4. 波普的证伪方法的局限性

波普的证伪方法对于现代科学来说是非常重要的，但是对这种方法也存在着大量的批评。批评者认为波普有关假说的论点可以夸大虚假的问题。尽管对黑天鹅的观察从逻辑上歪曲了所有的天鹅都是白色的一般概括，但是对真正的科学的假说却包含着更多的复杂性³⁰和或然性（非广泛性）。在新提出的假说依靠大量的被认为是真实的辅助假说这个层面上，它是复杂的。所以，如果从新的假说当中推理出来的预测在日后被证明是虚假的，就很难区分新的假说是错误的，还是在众多的辅助假说当中的某一个是不正确的。太阳系中的第八大行星——海王星的发现就恰好证明了复杂性这个情况。天王星的异常轨迹不是由于牛顿的运动定律本身的缺陷，而是由于太阳系只包含 7 大行星的辅助假说的缘故。然而，当这一理论在 20 世纪初期被证明是错误的时候，它确实被归结为牛顿的运动定律的缺陷。从事真正的科学是一件不容易的事。

另外，因为许多科学的假说存在一定的或然性，就像是在技术分析的领域一样，与假说相矛盾的观察可能从来不会被作为已经确定的虚假的证据。非常规的观察也许只是偶然发生的。为此，我们就要把统计分析引入进来。作为将在第 4 章和第 5 章讨论的内容，以被观察的证据为基础对一种假说进行反驳，会产生一定比例的错误。统计学会帮助我们确定这种概

① 1 码=0.914 4 米。

率的数量。

尽管还有许多其他的限制并不包含在我们的考虑范围之内，但是波普对科学的方法的发展所做出的贡献仍然是巨大的。

科学的假说的信息内容

简单回顾一下，如果一种假说能够做出可以检验的预测，那么它就是信息含有量丰富的假说。在这种情况下，就增加了这种假说被证明是错误的概率。因此，一种假说的可证伪性和它的信息内容之间就是相联系的了。

在具有科学意义的假设中，存在着信息内容和可证伪性的可信度问题。有些假说的信息含有量非常丰富，因此，它们也就比其他的假说更具有可被检验性。

波普所说的“大胆的猜测”是指从许多可以被检验的预测中推导出来的信息含量丰富的假说。这就是科学家的工作，所以，科学家们试图逐步反驳现存的假说，并且用信息含量更多的假说取代现有的假说。这种激励因素使得科学的知识得到了改进。

信息含量丰富的假说可以对范围广泛的现象做出精确的（小范围内的）预测。每一种预测都有机会证明这种假说是错误的。换句话说，一种假说的信息含有量越多，它被证明为虚假错误的假说的机会就越大。反过来，信息含有量低、受到约束的假说则很少会做出预测。这种情况的结果就是这些假说很难被证明是错误的。例如，某种技术分析法则声称在任何时间运用任何投资工具都会获得高额的利润，根据这句话我们就可以做出一种大胆的、包含有大量信息的判断，即可以通过其在某一时间段所进行的市場操作没有产生利润来证明这种法则的错误。相反，如果一种技术分析法则声称只能在一周的时间期限里对标准普尔的期货合约进行操作，并只获得微薄的利润的话，那么这个法则就是受到约束、信息内容含有量低下以及很难被检验的。对待这种情况，唯一反驳的方法就是对标准普尔 500 指数某一周根本没有获得利润的表现进行回溯测试。

有些技术分析方法看上去包含丰富的信息，而实际上它们却并不如此。艾略特波浪理论正好符合这一点。表面上看，艾略特波浪理论强有力地声

称在所有的时间范围和所有的市场中，价格的走势可以通过一种独特的、统一的原理进行描绘。这种论点似乎被艾略特波浪理论的实践者通过对以前的价格数据进行波浪走势的汇总所确认。而实际上，因为艾略特波浪理论没有对未来的价格走势做出过任何可以被检验的预测，所以，艾略特波浪理论在面对它是没有意义的评价时，就显得底气不足了。³¹

另一方面，让技术分析领域感到兴奋的是有效市场假说理论（Efficient Markets Hypothesis, EMH）的出现。在该理论中，术语“有效”指的是价格反映的是与资产未来收益率相关的所有已知和可知的信息的速度。在有效的市场当中，假设相关的信息会通过价格完全体现出来。有效市场假说分为三类。我们根据其受约束的程度、信息内容和可证伪性的降序来排列：强式有效市场假说、半强式有效市场假说以及弱式有效市场假说。

强式有效市场假说确认市场是有效的，并且充分反映了包括内幕信息在内的所有信息。强式有效市场假说认为，所有建立在有关收购的内幕消息之上的、以公开信息为基础的和以技术分析为基础的投资策略，是不可能获得跑赢市场的（超额的）收益率的。因为任何的投资策略，无论使用何种信息或分析形式所获得的异常收益，都会强有力地反驳有效市场假说，所以有效市场假说的这种大胆的说法看上去同样是包含有大量信息的，并且可以进行证伪的。然而，由于信息是无法证明是从私下获得的，所以这种说法在实际中是无法进行检验的。

半强式有效市场假说所包含的信息和范围就相对较少一些，也就是说市场仅仅在涉及所有已公开的信息的条件下才是有效的。半强式有效市场假说通过以公开的基础数据（市盈率、账面价值等）或者技术数据（相对强度等级、成交量比率等）为基础的投资策略得出跑赢市场的收益率，并以此对半强式有效市场假说进行证伪。实际上，半强式有效市场假说拒绝承认基础分析和技术分析的有效性。

弱式有效市场假说做出了大胆而又信息丰富的声明。它只反映了过去的价格、成交量和其他的技术资料。因为弱式有效市场假说只是拒绝承认技术分析的有效性，它给那些弄虚作假者提出了最困难的难题。他们唯一的愿望就是要提供以技术分析为基础的投资策略所产生的超额收

益率的实证。

由于弱式有效市场假说对于谎言来讲是最为不利的说法，因此，它也是最有可能被证明为是错误的，而对它的证伪同样会产生出人意料的结果。换句话说，在有效市场假说的三种观点中，对弱式有效市场假说的证伪可以使知识得到最大限度的增加。这种说法指出了科学的一般原理：当证明最受约束的和最为棘手的假说是错误的时候，我们就能够最大限度地获得知识。一项测试显示，来自公司总裁的内部消息是可以产生超额收益的（强式有效市场假说的证伪），它既不能给我们带来惊喜，也不能让我们从中受益。最重要的是，内部消息可以获得利润。除此之外，还有什么新鲜的吗？相反，对于技术分析的实践者和有效市场假说的支持者来说，对弱式有效市场假说的证伪就是件内容丰富的事情。它不仅意味着有效市场假说这条存在 40 多年的重要的金融学原理的消亡，而且还意味着对技术分析有效性的重要确认。两者都意味着现有知识的巨大变化。

因此，可以说对一种假说进行的证伪获得的知识，是与这种假说的信息内容成反比的。同样，也可以说根据对一种假说的确认（观察与预测相一致）所获得的知识与这种假说的信息内容有直接的关系。信息含量最丰富的假说对新知做出了最为大胆的声明。这些假说试图在用最少的假设的同时，用最为精确的方式理解最广泛的现象。当这种类型的假说被证明是虚假的时候，我们既不会感到吃惊，也不会从中学到任何东西。除非那些大胆科学家提出来，否则很少有假说能够被证明是正确的。例如，假设我们大胆地提出了一种新的物理学理论，其中的一种假设认为人们可以建造一种反重力装置。这件事情如果是真实的，那么这种理论就使得这一领域的知识得以创新。然而，如果这个装置不能够付诸实践的话，那么就没有人会对这个预测的失败感到吃惊了。然而，当一种保守的假说被证明是虚假的时候，它所预测的结果则是完全相反的。例如，保守的假说仅仅确信现在被人们接受的物理学理论是真实的，并预测反重力装置的建造是失败的。通过对反重力装置的工作情况进行观察来对这种没有说服力的假说进行证伪，会导致在知识领域的巨大收获——对新物理学的检验。

我们所提出的最保守的假说主张现在已经没有新的发现了。换句话说，

即我们现在知道的一切就是我们所知道的全部。这种假说拒绝承认任何主张有新的发现的假说的事实。主张没有新的发现的保守的假说在科学界有一个特殊的名称——**原假设**（null hypothesis），它对任何声称有新发现的主张的研究提出假设。不论主张确认一种新的疫苗可以治愈令人畏惧的疾病，还是告诉我们可以抵消重力的新物理学原理，或者是技术分析具有预测力，我们通常在初始的时候假设原假设是真实的。然后，如果实证证据不支持原假设的话，那么对于知识领域来说就可以获得巨大的收获。

因此，科学就会按照以下的步骤发展。当一种大胆的新假说被提出的时候，会产生反对的观点，即原假设。原假设和新的假说一样，都是大胆的假设。乔纳斯·索尔克（Jonas Salk）曾经有过一个大胆的假说，即他研制的疫苗可以预防小儿麻痹症，效果远比（试验药物用的）无效对照剂强得多，这个假说导致了矛盾的观点，即原假设。他做了一个保守的预测，即他的疫苗防治感染的效力不如（试验药物用的）无效对照剂。这个预测之所以会如此的保守，是因为这之前所研制的防治小儿麻痹症的疫苗全部失败了。这两种相互矛盾的观点只有成功与失败之分，没有第三种可能。如果其中的一种假说被证明是虚假的，那么按照排除中间选项的逻辑法则，我们会知道另外的一种假说是正确的。索尔克的实验证明了使用现存的疫苗比那些使用（试验药物用的）无效对照剂感染小儿麻痹症的比率要低很多。换句话说，这个实验充分证明了原假设预测的错误。这个使人吃惊的结果让医学界得到了巨大的收获。

科学家对证伪的反应

当经受了之前众多检验的假说或者理论由于最近的观察与预测出现矛盾并且被证伪时，科学家对此的反应又是怎样的呢？正常的反应应该是它能够为知识界带来最大的收获。因为科学家也是人，他们有时候也不能按照符合科学的方法做事。

对于知识的收获，有两种可能的反应。第一种是通过利用这些知识去预测新的、未知的事实，来保留现存的假说。如果这些新的事实能够得到确认，并且可以解释与假说相矛盾的观察资料与假说不再矛盾的原因。然

后, 这个假说就会得以保留下来。新的事实代表了我们对世界已知的内容得到了充实。第二种反应就是拒绝接受旧的假说, 而提出新的假设, 新的假设不仅对所有经过以前的假说解释过的观察资料做出满意的解释, 而且也可以对新产生的不一致的观察资料做出解释。这种方式也代表了新的知识内容得到了充实, 只不过它是通过以更加深入的探索和预测力的形式得到的。然而, 不论发生哪种情况, 正确的反应就是做任何能够促进知识前沿发展的事情。

不幸的是, 科学的兴趣有时候会让步于个人的目的。人之本性会阻碍优秀的科学的发展。以前的假说在情感上、经济上以及专业上的束缚, 可以在某种程度上激发对减少假说的信息内容和可证伪性的不一致的证据的辩解。这种辩解使得知识的前沿边界发生了倒退。幸运的是, 科学本身是可以进行自动修正的整体。当科学家的研究偏离了正当的轨道时, 科学家这个群体会非常愿意指责这些误入歧途的“兄弟”。

有些样本可以证明这些抽象的概念。首先, 我提供这样一个样本, 即通过新的观察资料对与预测相矛盾的事情做出适当的反应。在这个案例当中, 新的事实试图对以虚假为基础建立起来的理论进行辩解。这种方式是19世纪的两位天文学家所采用的, 他们通过发现新的行星的方式, 导致了新知的出现。在当时, 牛顿的运动定律和万有引力定理是已经被人们所接受的关于行星运动的物理学。这两个定律通过无数次的观察被反复地确认着, 但是出乎天文学家意料的是, 某一天, 新型的、高倍的望远镜显示天王星正在偏离运行轨道, 这一现象对牛顿的运动定律提出了质疑。对证伪的刻板、不恰当的应用要求立刻抛弃牛顿力学。然而, 所有的理论都是以辅助的假设为基础的。这个案例中的一个最重要的假设就是, 太阳系中只包括7大行星, 而天王星是第7个也是离太阳最远的一颗行星。这个假设导致了天文学家亚当斯(Adams)和列维叶(Leverrier)大胆地预测还没有被发现的第8颗行星(新的事实)的距离, 它的距离比天王星还要远。如果这个假设是真实的, 那么新的行星的引力影响就可以解释天王星的异常运行轨迹了, 很显然, 这个理论与牛顿的力学理论相矛盾。另外, 如果第8大行星确实存在的话, 牛顿的两个定律之中的一个就能够预测到可以被

观察到的新的行星在天空中的具体位置。这是一个大胆的、信息含有大量的以及可以进行高度检验的预测，它使得牛顿的定律受到了最严格的检验。随着海王星的发现，天文学家亚当斯和列维叶对于新的行星会出现在天空中的大胆预测在 1846 年得到了证实。他们通过证明牛顿的定律不仅能够解释天王星的偏离运行，而且还可以做出高度精确的预测，进而使牛顿的定律得以从质疑中保留下来。正是这种适当的方法使得一种理论在面对自相矛盾的观察资料时，仍然可以保留下来。

然而，这个案例在牛顿所有的定律和理论面前，只能证明牛顿的定律只是暂时正确的。尽管牛顿的定律已经被人们毫无疑问地运用了 200 多年，但是在 20 世纪初期已经有更多精确的天文观察被证明与牛顿理论所做的预测不符。旧有的理论最终会被证明是错误的，而新的理论将取而代之。1921 年，阿尔伯特·爱因斯坦通过提出了他的新的、含有更多信息的广义相对论对这种观点做出了适当的回应。在将近 100 年以后的今天，爱因斯坦的理论在众多的质疑声中存活了下来。

牛顿的理论之所以会被称为科学是因为它为实证的反驳提供了空间。实际上，牛顿的理论并不能够算作错误，但至少是不完整的。爱因斯坦提出的广义相对论不仅对牛顿模型覆盖的所有现象进行了说明，而且对新的观察资料逐步适应，而这些观察资料是与牛顿理论的局限性相矛盾的。科学就是这样不断进步的，而一种理论的寿命长短与其所处的地位没有任何意义。

天文学家亚当斯和列维叶的案例清楚地解释了一种观点必须为实证的反驳提供空间的原因。单独的证伪使得科学的知识比传统的观点更具优势。抛弃错误的或者不完整的思想，并且用信息含量更大的观点取代旧有的思想的能力可以建立一种能够进行自我修正的知识体系，并且将这种状态一直保持下去。这种能力反过来也为新的思想建立并达到更高的理解水平提供了坚实的基础。从事不能消除错误的知识程序的脑力劳动，会不可避免地陷入荒谬的境地。这一点也是流行的技术分析所面临的问题。

现在我们来考虑一个对虚假的事物做出不适当反应的案例。这个案例发生在金融领域。将科学引入金融领域是最近才有的事儿。也许这样可

以对有效市场假说的支持者所做的辩解和非科学的反应做出解释。当观察资料与自己所信奉的理论相矛盾的时候，他们会通过减少这种理论所包含的信息内容使得该理论得以保留下来。正如我们之前提到的包含最少信息的弱式有效市场假说理论，它是以技术分析为基础的投资策略做出预测的。³² 这种形式的有效市场假说是不可能获得跑赢市场指数的经过调整的收益率的。当有效市场假说的支持者面对基于技术分析的策略仍然可以获得超额收益的事实时，³³ 他们会通过将自己持有的理论内容最小化来对此做出回应。有效市场假说的支持者会加入新的风险因素并且宣称依靠技术分析获得的超额收益仅仅是对这种策略固有的风险做出修正的结果。换句话说，有效市场假说的辩护者声称坚持技术分析策略的投资者会将他们自己暴露在他们所特有的策略的风险之下。有效市场假说的支持者因此会将技术分析策略的收益描述为非跑赢市场的收益。现在回想一下有效市场假说没有声称获得跑赢市场的收益是不可能的这句话吧。它仅仅是做出了获得高收益必须承受额外风险的假设罢了。如果通过技术分析策略获得的收益确实要对更高的风险做出补偿的话，那么有效市场假说就会得到完整的保留，尽管研究显示技术分析策略所获得的收益率要高于市场指数。

有效市场假说的支持者这样做会产生一个问题。在技术分析研究已经进行之后，他们会编造风险因素。³⁴ 这不能理解为正确的科学。如果有效市场假说在进行研究之前就否认了风险因素并且预测对技术分析的检验会产生类似跑赢市场的收益率，那么从科学的角度来看，它就是正确的。他们在做了这些工作之后，有效市场假说的地位是否会因为成功的预测而得到加强呢？相反，当有效市场假说的支持者需要随时对与自己所持有的理论相矛盾的发现进行辩解的时候，他们会采取欺骗或者不道德的手段虚构一个新的风险因素来保护他们的假说。只有这样，有效市场假说的支持者才会把这个假说包含的任何信息内容都利用起来，最终放弃不能被证明为不正确的假说。

这种有效市场假说对证伪做出不敏感处理的知识回归方法先例，已经通过对以前的有效市场假说进行有缺陷的辩护建立了。这些早期的、带有误导性的对有效市场假说进行辩护的努力，对市净率、市盈率等公开的

基本信息可以产生超额收益率³⁵的研究做出了回应。对于这种有缺陷的实证的回应，有效市场假说的辩护者声称低市净率和低市盈率仅仅是高风险市场的标志。换句话说，相对于其自身的账面价值而言，过低的股票价格暗示了该公司目前正处于困境之中。当然，这种推理是要经过循环论证的。目前最主要的问题是在研究表明具有这些特征的市场可以产生超额收益率之前，有效市场假说没有将低市净率和低市盈率定义为风险因素。一旦有效市场假说的理论家这么做了，那么他们就会用额外的风险尺度来支持有效市场假说的新内容。相反地，有效市场假说的理论家出于对已经发生的不一致的观察资料的辩解的具体目的，在事后发明了这些风险因素。像这种解释我们称为特设假设（ad-hoc hypotheses），即出于对一种正在进行证伪的理论或者假说保护的目的，在事后所做出的解释。波普把这种知识回归，以及不惜一切代价保护某种理论的行为称作证伪免疫（falsification immunization）。

如果波普发现了这种情况，他就会对有效市场假说的铁杆支持者提出批评，但是他可能会对那些提倡行为金融学的人所做的努力而喝彩。行为金融学这一相对新兴的领域提出了可以进行检验的假说，即以源于认知偏差和发明者的幻觉的公共技术和基础数据为基础的策略的收益率进行解释。具有讽刺意味的是，对主观的技术分析的有效性和某些形式的客观技术分析盈利的有效性的错误的信仰，都有可能是认知缺陷导致的结果。

科学的态度：开放与质疑并存

证伪主义在科学发现的两个方面做出了明确的区分：提议与反证。这两个方面需要不同的思维模式——开放性和质疑性。两种不同思维模式的共存对科学的态度给出了定义。

当假说已经成立的时候，对一种新思想抱有开放的态度是十分重要的。希望用新的方式去了解新事物、获取新的解释，以及做出大胆的归纳的愿望³⁶构成了提议方面的特点。大多数的技术分析从业者在这一点上都做得非常出色。而新的指标、新的体系以及新的形态随时都会被提出。

然而，一旦一种大胆的假设形成，它的可接受性就要由提议向质疑的

方向转变。对新思想的怀疑促使了对其缺陷进行严格、连续的搜索。因此，在科学家的大脑中，他们对猜测性的好奇与不懈的质疑之间一直保持着紧张的神经。然而，怀疑并不是指坚持不懈的质疑，而是指向有说服力的新实证进行妥协。正是这种准精神分裂的状态对科学的态度做出了定义。

除了对新假说的质疑，另外一种形式的质疑在科学家的思想里萌发：对这种思想自身进行质疑。这种想法来源于对一种人性倾向的深刻认识，即人性通过仓促的概括，从而导致了无事实依据的结论（参见第2章）出现的倾向。科学的这一过程可以视为反对这些倾向性的保障。

最终结果：假设演绎法

有些人也许会说根本就没有什么科学的方法。³⁷ “作为一种方法，科学的方法只不过是某人没有被束缚的大脑尽最大的努力所做的事情而已。”³⁸ 科学的方法的本质就是解决智能问题。

今天在科学领域里所应用的解决问题的方法被称为假设演绎法。它通常分为5个步骤，即观察、假设、预测、检验和结论。“实际上，在科学的工作中，这5个步骤是交织在一起的，这对于把任何科学调查的历史进行严谨的分类来说，无疑是件困难重重的事儿。有时候不同的步骤会混合在一起或者彼此间的关系模糊不清，而且它们通常不会在结论清单上显示出来。”³⁹ 另一方面，这种方法可以帮助我们对过程进行思考。

假设演绎法是牛顿于17世纪首先提出的，但是这一理论直到波普的论文发表之后才正式得到命名。假设演绎法是科学家和哲学家几百年来不断争论的自然结果。假设演绎法把归纳法与演绎法逻辑结合在一起，在平衡各自力量的同时，也对各自的局限性加以限制。

5个步骤

(1) 观察。一个潜在的形态或者联系，在之前的一组观察资料中被注意到。

(2) 假设。基于无法解释的洞察力，以前的知识以及归纳概括的结合，

我们假设看到的形态并不是一组特定的观察资料中人为创造的，而是应当出现在类似观察资料中的实际形态。这个假设仅仅对这个形态是真实的做出肯定（科学的法则），或者进一步对这个形态存在（科学的理论）的原因给出解释。

（3）**预测**。预测是指对假说和包含在有条件的命题中的假设进行推理。命题的先行从句是假设，结果从句是预测。如果这个假设是真实的，那么通过预测我们可以知道在新的观察资料组里面应当如何选择观察的对象。例如：如果某个假设是真实的，那么当运行 O 进行的时候，X 就会被观察。由 X 所定义的一组结果就会清楚地告诉我们哪个将要进行的观察是与预测相吻合的，而且更重要的是，我们还可以知道哪个观察是与预测相矛盾的。

（4）**检验**。新的观察将与特殊的操作保持一致，并且与预测进行比较。在某些科学领域中，这种操作实际上就是受到操控的实验。而在其他科学领域中，这种操作是指观察研究。

（5）**结论**。对于假说是真理还是谎言的推理是以哪种观察资料与预测的确认程度最高为基础的。这个步骤包括诸如置信区间、假设检验在内的统计学推理方法。这一部分我们将在第 4 章和第 5 章进行描述。

来自于技术分析的样本

假设演绎法将应用于有关技术分析的新思想的检验当中，以下就是假设演绎法的样本。

（1）**观察**。当一个股票市场的指数，例如，道琼斯平均指数或者标准普尔 500 指数站上 200 日移动平均线的时候，一般说来指数还会在未来的几个月当中继续上涨（概率泛化）。

（2）**假设**。在对技术分析进行的观察、归纳概括以及早期发现的基础上，我们提出下列假设：当道琼斯工业平均指数（DJIA）向上穿过 200 日移动平均线时，一般情况下，多头在未来 3 个月是可以赚得利润的。我对此种假设是以 200 小时为标准的。

（3）**预测**。在假设的基础上，我们预测一项观察研究，或者说是回溯测试，会产生利润。假设和预测将转入下列有条件的陈述当中，即如果道

琼斯工业平均指数向上穿过 200 小时移动平均线是真实的，那么进行回溯测试就可以产生盈利。然而这个预测产生了一个逻辑问题。即使回溯测试是可以产生盈利的，那它对于证明指数向上穿过 200 小时移动平均线就是正确的也没有太大的帮助。因为就像我们之前指出的那样，当可以产生盈利的回溯测试与指数向上穿过 200 小时移动平均线的正确叙述是相一致的时候，它是不能证明指数向上穿过 200 小时移动平均线就是真实的。如果我们试图这样做的话，就会承认对所确认的结果是错误的（参见辩论 1）。另一方面，如果回溯测试显示的结果是不能够产生盈利的话，那么用有效的逻辑方式对结果进行证伪就会有力地证明指数向上穿过 200 小时移动平均线的假设是错误的（参见辩论 2）。

辩论 1

假设 1：如果 200 小时的假设是真实的，那么进行回溯测试就会产生盈利。

假设 2：回溯测试确实是产生盈利的。

无效的结果：因此，200 小时的假设是真实的（对结果的确认产生的错误推论）。

辩论 2

假设 1：如果 200 小时的假设是真实的，那么进行回溯测试就会产生盈利。

假设 2：回溯测试没有产生盈利。

有效的结果：因此，200 小时的假设是错误的。

然而，我们的目标就是要证明 200 小时的假设是正确的。为了解决这个逻辑问题，假设我们在第 2 个步骤中已经建立了一个原假设，即当道琼斯工业平均指数（DJIA）向上穿过 200 日移动平均线时，在未来的 3 个月当中并没有产生盈利。我们将其简化为原 200（Null-200）。通过这个假设我们可以成立下列有条件的命题：**如果原 200 的假设是真实的，那么进行回溯测试将不会产生盈利。如果回溯测试无法证明这个假设是可以盈利的，我们就可以说这个原假设是错误的（错误的结果）。根据排中原则，不**

论 200 小时的假设与原 200 的假设哪个是真实的，都不会出现第 3 种答案；也不存在其他假说的可能。因此，通过反驳原假设，我们将间接证明 200 小时的假设是正确的。我们来看下面这个有条件的三段论（演绎推理）。

假设 1：如果原 200 的假设是真实的，那么进行回溯测试将不会产生盈利。

假设 2：回溯测试确实是产生盈利的。

有效的结果：原 200 的假设是错误的，所以 200 小时的假设是真实的。

（4）检验。提出的法则经过回溯测试，它的收益率便可以观察到。

（5）结论。确定结果的意义是统计学推理的工作，我们将在接下来的 3 章中进行讨论。

对观察结果进行严谨和详细的分析

假设演绎法的第 5 个步骤指出了科学与非科学之间存在的另外一个重要的差异。在科学的领域中对证据的观察并不是来自事物表面上的意思。换句话说，实证看上去最为明显的含义也许并不能够代表其真正的含义。在得出最后的结论之前，实证都要接受最严格的分析。科学领域中对于实证的选择是一个定量数据，而对用于推导结论的工具的选择则依靠统计推断。

一种最重要的科学原理就是对于简单的解释（奥卡姆剃刀）的偏爱。正如其名所示，只有在普通的假说被排斥之后，惊人的假说才会被考虑进来。人们目睹 UFO 的出现并不能马上作为外星人访问地球的实证。世俗对球状闪电、气象气球或者是新型飞行器的描述在开始的时候是不能相信的，除非我们遭到了来自外层空间的侵略。

因此，我们首先需要用科学的态度看待对能够获得丰厚利润的法则进行的回溯测试，并且在考虑到一种有效的技术分析法则被发现的概率之前，要拒绝其他的解释。当我们取得的良好业绩与一种法则的预测力没有任何关系的时候，我们最可能对此做出的解释就是由于样本的错误（参见第

4~5章)所导致的好运气,以及由于数据挖掘(参见第6章)产生的系统性错误。

对科学的方法的主要内容总结

下面所列出的就是科学的方法的主要观点。

- ▶ 不论实证的数量有多少,科学的方法从来没有对一种假说的真实性做出令人信服的证明。
- ▶ 被观察的实证和以演绎方式对结果进行的推理相结合,可以用来证明有具体的成功概率的假说是虚假的。
- ▶ 对于可以检验的假说来讲,科学是有其局限性的。不能够被检验的论点是不属于科学讨论的范畴的,而且其本身也不具有任何意义。
- ▶ 如果,并且仅仅如果对于尚未进行观察的资料的预测能够从某种假说当中推理出来的话,那么这个假说就是可以进行检验的。
- ▶ 一种能够对过去的观察资料进行解释,但是不能够对新的观察资料进行预测的假说,是不科学的。
- ▶ 我们通过对新观察资料和假说所做的预测进行比较的方式,对某种假说进行测试。如果预测与观察资料的表现相一致,那么这种假说就是没有经过证明,而仅仅是得到了暂时性的确认。反之,如果预测与观察资料的表现不相一致,那么这个假说就被证明是虚假或不完整的。
- ▶ 当下所有被我们接受的科学的知识的正确性只是暂时的。在将其证明为虚假的知识之前,它仍保持着真理的地位,而当其的虚假性得到确认的时候,它就会被更加完整的理论所替代或者涵盖。
- ▶ 科学的知识是不断积累并且渐进地向前发展的。当旧有的假说被证明是错误的时候,它们就会被能够明确反映客观现实的新假说所替代。对于研究或者声称新假说是更好的知识学科来说,科学就是唯一的方法。虽然音乐、艺术、哲学或者文学等其他学科的知识、风格和方法可以随着时间的推移而改变,但是这种改变不能说新的理

论就比原来的理论要好。

- 任何一组过去的观察资料（数据）都可以通过无限的假说进行解释。例如，艾略特波浪原理、江恩线、经典的图表形态以及其他方法都可以根据自己的分析方法对过去的市场行为做出解释。因此，所有这些解释都可以被认为等同于实证。至于决定哪个方法是更好的，唯一的方法就是看看它们对还没有得出结果的观察资料所做的预测吧。那些对新的观察资料不能够产生检验性（证伪）的预测将会被立刻淘汰，因为它们的存在是毫无意义的。而那些可以产生预测，但又与未来的观察相矛盾的方法也会由于它们被证明是错误的而被淘汰。所以，只有能够显示出真正的预测力，以及能够进行检验的预测方法才会保留在技术分析的知识体系中。

如果技术分析采用科学的方法

在这一部分中，我们将要分析采用科学的方法的技术分析的结果。

淘汰主观的技术分析法

采用科学的方法的主观的技术分析最重要的结果就是将主观的方法淘汰。因为它们是不可检验的，主观的方法可以经受得住实证的挑战。这就使得它们比错误本身还要糟糕。它们是缺少信息的没有意义的论点。主观方法的淘汰使得技术分析面临着彻底的客观实践。

主观的技术分析的淘汰主要通过两个方面：转变为客观的方法或者被抛弃。也许江恩线、主观的分歧、趋势通道以及其他众多的主观形态和概念包含着有效的市场行为特征。然而，对于它们提供的主观形式，我们是拒绝接受的。

将主观的方法转变为客观的方法是没有意义的。为了显示一个这方面的实例，我将对头肩形态的自动检测和“主观的技术分析的客观化的测试结果”这一运算法则进行讨论。

淘汰没有意义的预测

假设所有的主观的技术分析从业者都按照我所呼吁的那样将主观的技术分析客观化或者全部放弃是不切实际的。对于那些仍然使用主观方法的人而言，如果他们的方法论不是客观的话，有一种重要的途径可以让他们的成果被接受采用。从此以后，他们可以提出可以被检验的预测。这种方法可以使得他们所提供的信息含量丰富而又具有意义。在这种情况下，他们所提供的信息并不需要是正确无误的，而是他们的预测中含有认知的内容，可以通过我们在导论中讨论过的可以识别的差别测试。换句话说，预测可以总结部分内容，并通过市场的后续表现判断哪个预测是正确的。或者说，预测将清楚地暗示哪个结果是错误的。正如我们之前所指出的，从本质上来讲，一个不能做出哪种未来事件能够导致错误的预测，可以产生任何结果。

现在大多数主观的预测，正如市场所说的那样，是没有意义的。在有可能性当中，是主观技术的实践者还是分析师做出的预测并不明显。首先，我们考虑一个明显没有意义的预测：“我所应用的指标预测市场要么上涨的百分比是没有限制的，但下跌的幅度是 100%，要么涨跌幅度在两者之间。”从表面上看，这段陈述是可以被检验的，因为没有任何一个明确的结果与所做出的预测相矛盾。对于这个预测来说，唯一合理的事情就是预测的意义及其可证伪性的缺失是如此的明显。一个更加典型的市场如下面所示：“根据我的指标（加入一种或者多种技术分析方法），我对市场是持牛市观点的。”像这种不能够进行证伪的表述就是没有意义的，但是它在内容上的缺失是不明显的。尽管它做出了上涨的预测，但是具体到什么时间、在哪种条件下预测会出现错误就表述的不是那么清楚了。

我们可以通过排除若干种结果的方式，使这种牛市的观点变得富有含义。例如，我期望市场在基于现有的水平下跌超过 5% 以前，能够在当前这个水平的基础上涨超过 10%。任何在上涨 10% 以前下跌超过 5% 的观点都会被归为错误的预测。

如果你曾经怀疑自己正在提出一个没有意义的预测，那么现在正好有一些好的补救方法。请提出下列问题：“在你承认这是一个错误的预测之时，

相反的走势（与正确的预测相反）已经发展到什么程度了？”或者“你的预测排除了哪些结果？”再或者“何时，并且在什么情况下，对诸如时间的推移、价格的改变（相反的或者正确的）和特殊指标的形成之类的预测进行检验？”各位读者，如果你对看到人们的局促不安感到过敏的话，你最好不要这么做。

为了使主观的预测具有意义，对于预测在什么时候、进行怎样的评估做出事先的说明可以不给市场权威人士做事后诸葛亮留有任何的余地。给主观的预测添加含义的方法包括以下几点：①当要对预测做出评估的时候，给未来的某个时间点做出定义；②在不宣布预测是错误的前提下，对能够允许相反走势偏离的最大程度做出定义；③在预测具体的利空的运行幅度之前，先预测具体的利多的运行幅度（在 Y% 的利空运行幅度之前做出 X% 的利多运行幅度）。像这样的步骤将会允许主观的实践者制作一个富有含义的市场轨迹记录。有意义的轨迹记录可以从实时进行的具体交易建议中获得。

这种建议的局限性之一是它会导致可以盈利的轨迹记录所表述的内容模糊不清。因为主观的预测是通过模糊的方式推导而出的，即使一个包含丰富预测内涵且能够盈利的轨迹记录可以随着时间的推移建立起来，但它也不能够被认为是可以在未来进行重复的、并且做出一致预测的结果。实际上，随着时间的推移，分析方法也不一定一直保持稳定。对专业的主观判断的研究表明专家们不能将在一种判断中得到的一系列信息与接下来的信息相结合。“直觉通常会受到一系列随机出现的由疲劳、厌倦以及所有影响我们人类⁴⁰的矛盾的困扰。”换句话说，在不同的时间给出相同的市场数据形态，主观的技术分析师很有可能会做出不同的预测。⁴¹

模式变化

将技术分析再次转变为客观的实践，指的就是托马斯·库恩（Thomas Kuhn）所说的模式变化。在他出版的具有高度影响力的著作《科学革命的结构》（*The Structure of Scientific Revolutions*）一书中，库恩否认了人们持有的“科学包括严格的证伪和推理”这一概念。他认为科学的进步是由一系

列模式和世界观组成的。当一种固定的模式确立之后，受到这种观点引导的实践者就会限制他们自己提出与这种观点相一致的、并且可以用这种观点进行解释的问题以及进行假设的活动。

大量的技术分析师被灌输了非科学的知识，直觉分析的模式是由技术分析的前辈们，比如：道、江恩、沙巴克、艾略特、爱德华兹和迈吉等发展壮大的。他们建立了主观的研究传统，并且对其背景知识进行了假设，然后再将这些被认为是正确的知识传授给那些热情的实践者。由美国纽约市场技术分析师协会（MTA）举办的特许市场技术分析师（CMT）资格认证考试就是很好的范例。

向实证客观方法的转变将会对既没有意义，又没有充分的统计学实证支持的理论发起挑战。很多的学说将会重蹈古希腊物理学和天文学的覆辙。有些方法会经受客观和统计学的检验而保留下来，并且在正式的技术分析知识体系中占有一席之地。

也许读者会认为我的观点十分尖锐，我并不提倡证伪的标准适用于消除所有正在形成的技术分析方法。很多杰出的科学理论都是开始于近代科学以前的不成熟思想，并且是以错误的证伪标准为基础的。这些思想需要时间去发展壮大，终有一天它们会成为富有内涵的科学。作为艾略特波浪原理的产物，社会从众心理驱动下的宏观经济表现（socionomics）这个领域迎来了它的新成员，即技术分析当中的样本。现在，我把这种新近发展的理论称为“近代科学以前的理论”，尽管它有着成为科学的潜质。根据我与约翰·诺夫辛格教授（从事社会从众心理驱动下的宏观经济表现领域的工作）的交谈记录，此时这种原理还不足以做出可以进行检验的预测。这就需要对社会情绪进行量化的设定，并根据社会从众心理驱动下的宏观经济表现对市场的走势做出判断。

初期的研究范围不应该仅仅因为不能产生可以被检验的预测而被缩小。我相信他们有朝一日会这么做的。

主观的技术分析客观化：样本

推动技术分析向科学方向发展所面临的挑战之一就是主观的图形形

态转变为客观定义的, 可以进行检验的形态。这一部分给出了两种专业技术对技术分析进行转变的样本——即 Keving Chang 和 Carol Osler (C&O) 对头肩形态进行客观化转变的样本。⁴² 并不是所有的形态内容都包括在内。在此我尽量列出他们所持理论的内容, 以及他们所面临并解决的问题, 并且举例说明将主观的方法转变为客观的方法所要面对的挑战。具体的细节请参见他们两人的原始资料。

我们在很多关于技术分析的文献⁴³当中都可以看到对头肩形态的描述, 这种形态表现为一系列无噪声的价格波动, 图形基本上与图 3-8 相类似。当这种形态的出现与教科书上所示的内容完全一致时, 即使是刚刚入门学习技术分析的学生都可以马上把它认出来。

那么当实际的图形形态与标准的头肩形态不相符合时, 问题就出现了。即使是经验丰富的图形分析专家也拿不准眼前的这个图形是否就是我们已经熟知的头肩形态, 进而缺少客观形态定义的问题就不可避免地发生了。主观的形态定义一般情况下只是对头肩形态大概的模样进行了描述, 但是他们并没有提供一个明确的将符合头肩形态的图形从那些只有一部分符合, 但是又不完全具有头肩形态特征的图形当中区分出来的法则。换句话说, 缺少明确法则定义的图形应当被排除在外。关于这个问题的概念如图 3-6 所示。

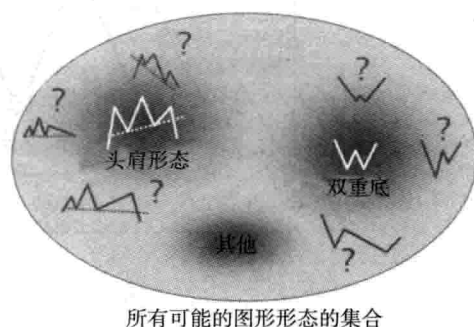


图 3-6 主观的形态——没有明确的排除规则

在不能够用客观的法则判断究竟哪些形态符合, 哪些形态不符合公认的头肩形态的标准时, 我们是不可能对这种形态的盈利性和预测力做出评估的。解决这个问题的方法就是用客观定义的法则将有效的头肩形态从无效的头肩形态当中分辨出来。⁴⁴ 这个概念我们在图 3-7 中进行列举。我们可以把主观的形态转变为客观的形态所面临的挑战, 理解成为对存在于所有技术分析价格形态超集其中的一个明确的子集进行定义的问题。

C&O 将头肩形态的顶部形态分为了 5 个支点，或者说是价格反转的 5 个支点，反过来，即 3 个最高点和 2 个低点。我们在图 3-8 中通过对 A 点到 E 点的标注来表示。在所有的 8 个点当中，C 点标注的是最高点，它要比围绕在它周围的两个最高点（即 A 点和 E 点，我们称为肩）还要高。目前在图形分析师中还对诸如“双下巴”、“喉结”以及“螺旋状头发”等具有辅助性特点的形态存在着争论。然而，由 C&O 定义的技术分析却与这些概念不一致，而且它们也不包括在里面。

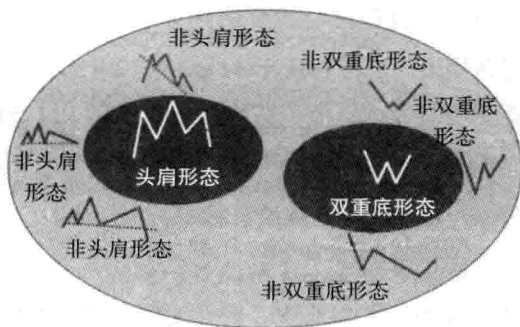


图 3-7 客观的形态——明确的排除规则

高。目前在图形分析师中还对诸如“双下巴”、“喉结”以及“螺旋状头发”等具有辅助性特点的形态存在着争论。然而，由 C&O 定义的技术分析却与这些概念不一致，而且它们也不包括在里面。

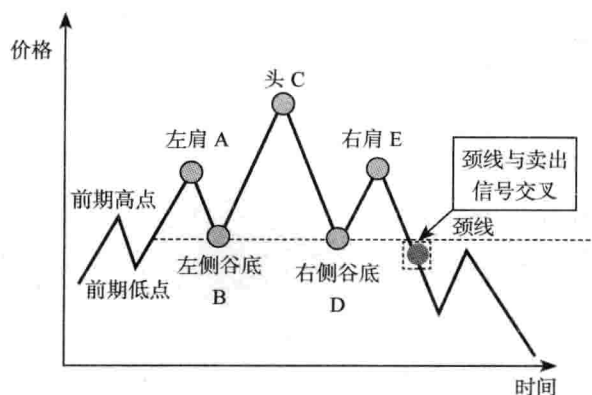


图 3-8 头肩形态

现实中的图形分析师所面临的一项挑战就是实际的价格波动并没有描绘出清晰的、可以辨认的高点和低点。相反，高点和低点出现在了巨大的时间范围内，从持续时间最短的数分钟到最长的几年到几十年不等。我们把这种特性称为不规则的范围，这对于主观的分析师试图确认头肩形态来说，无疑加大了难度。分析师必须在一个特定的范围内通过肉眼将价格行为先过滤掉，然后再单独考虑高点和低点。这种做法在进行事后的回顾时相对容易做到，但实际上，在图形的走势还没有完全明朗之前，是非常难

判断头肩形态的。

C&O 通过“亚历山大过滤器”⁴⁵或者“Z 字形指标”这类百分比过滤器，对这个问题进行了强调。“亚历山大过滤器”这一概念在梅里尔 (Merrill) 的《过滤波浪》(Filtered Waves)⁴⁶一书中进行了讨论。对于在价格的变化过程中确认高点和低点来说，这是一种客观的方法。当价格从最近的低点上涨到超过一个特定的百分比临界点，然后价格再从最近的高点下跌到那个特定的百分比临界点时，高点便可以确认了。例如，如果这个临界值设定为 5%，而价格却在最近最高价格的基础上下跌了至少 5%，此时的高点就没有出现。同理，如果价格在最近的价格最低点上涨了至少 5%，此时的低点也没有出现。对于价格最小波动的设定会导致高点和低点实际出现的时间与通过“Z 字形指标”过滤器得出的时间之间出现滞后的现象。

接下来，C&O 强调了如何确定正确的临界百分比，来对“Z 字形指标”进行定义。不同的过滤临界值会反映出不同百分比数量（范围）的头肩形态的。例如，一个 3% 的过滤器临界值所反应的头肩形态，相对于用 10% 的过滤临界值所反映的头肩形态来说，完全可以忽略不计。这就使得在不同范围内的多重头肩形态得以共存。C&O 强调这个问题的时候，是通过将每一种金融工具（股票或者货币）代入 10 种不同的“Z 字形指标”之后获得的一系列临界值的方式来体现的。此举使得他们在各种不同的范围内确认头肩形态成为了可能。

然而，这又引起了另外一个问题。10 个过滤临界值的含义是什么？显然，一组临界值也许对一种金融工具有效，但是对其他类型的金融工具来说就未必有用，因为每一种金融工具波动的程度各不相同。意识到这一点，C&O 把每一种金融工具最近的波动情况考虑进去，从而得出了一组 10 个过滤临界值。这样就使得这些金融工具的头肩形态计算法则可以对市场的各种波动进行概括。C&O 把市场的波动性定义为 V ， V 代表作为最近 100 个交易日每日价格变动百分比的标准差。10 个临界值是通过用 V 乘以 10 个不同的系数得出的，即 1.5、2.0、2.5、3.0、3.5、4.0、4.5、5.0、5.5、6.0。由此得出的 10 个“Z 字形过滤值”代表着不同的敏感程度。由 C&O 选择

的这组系数的波动性，是通过技术分析的从业者对价格图形进行肉眼检验得到确认的。而这些从业者是认同这 10 个“Z 字形过滤值”在确认头肩形态时起到了有效的作用的。

C&O 随后讨论了关于对有效的头肩形态进行规则定义的问题。这些法则应用于金融工具当中，一旦这种金融工具的价格出现了“Z 字形过滤”的现象，那么在给定的范围内高点和低点就会立刻得到确认。

首先，根据 C&O 的运算法则确认的头肩形态必须符合以下两个基本条件。

- (1) 头部的形态必须高于在其左边和右边的两个肩部。
- (2) 这种金融工具在头肩形态形成之前，必须是上涨的走势。因此，该形态的左肩就必须高于之前的最高点（PP），而左边的低点也必须高于之前的低点（PT）。

其次，C&O 还努力解决更多的细节和复杂的事情，从而使得有效的头肩形态符合头肩形态的要求。他们用一组创新的测量法达成了这一目的，即用垂直水平的对称图形来符合头肩形态的标准，并且依靠时间来将这个图形最终完成。这些法则使得他们可以把既是头肩形态，又不是头肩形态的形态最终确定为头肩形态。

垂直对称法则

垂直对称法则不包括颈线斜率过大的图形。如图 3-9 所示的图形是符合垂直对称法则的。

本法则用 A 点与 B 点的中间点 X，与 D 点和 E 点的中间点 Y，对右边的价格水平与左肩（A 点和 E 点）以及左边的低点（B 点和 D 点）进行了比较。为了确认垂直对称的头肩形态，这一形态必须符合如下的规则。

- (1) 左肩的价格水平线的高点 A，必须超过 Y 点的价格水平线。
- (2) 右肩的价格水平线的高点 E，必须超过 X 点的价格水平线。
- (3) 左边价格水平线的低点 B，必须低于价格水平线 Y。
- (4) 右边价格水平线的低点 D，必须低于价格水平线 X。

图 3-10 和图 3-11 显示的两种头肩形态将会被排除在外，因为它们不

符合垂直对称法则的要求。

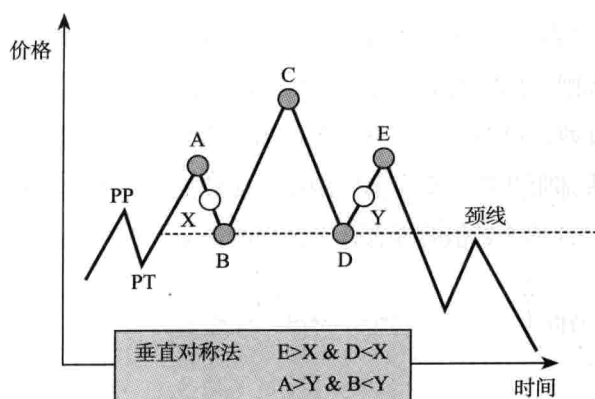


图 3-9 标准的垂直对称头肩形态

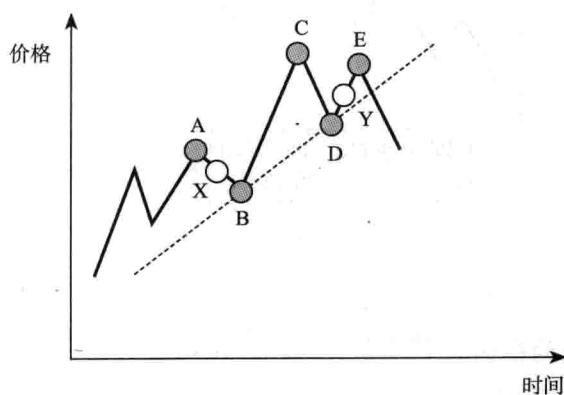


图 3-10 不合格的垂直对称头肩形态——颈线斜率过大

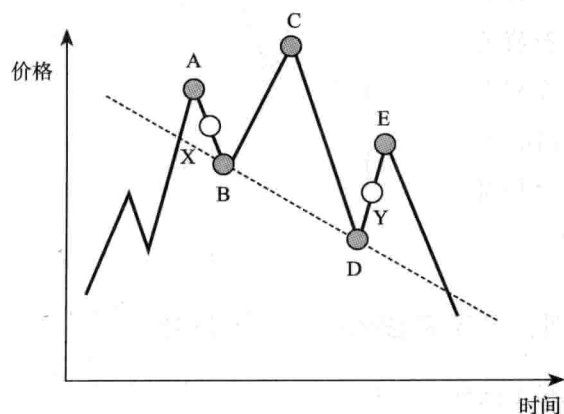


图 3-11 不合格的垂直对称头肩形态——颈线斜率过大

横向对称法则

C&O 用于从非头肩形态当中分辨出正确的头肩形态的另外一种方法就是横向对称法则。标准的横向对称是指某种形态的头部，即图 3-12 中的 C 点，与代表其两个肩部（A 点和 E 点）的高点的距离大致相等。C&O 的这一法则是说头部距其中一个肩部的距离不应该大于头部到另外一个肩部距离的 2.5 倍。2.5 这个数值并没有什么奇特的含义（见图 3-12）。

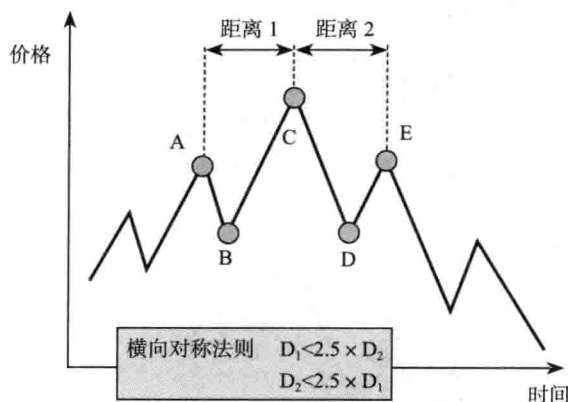


图 3-12 标准的横向对称

如图 3-13 所示的是不符合标准的横向对称法则的图形形态。请注意右肩的延伸长度过大。凡是右肩的延伸长度大于左肩距离头部长度的图形，都被视为不符合标准的横向对称图形。

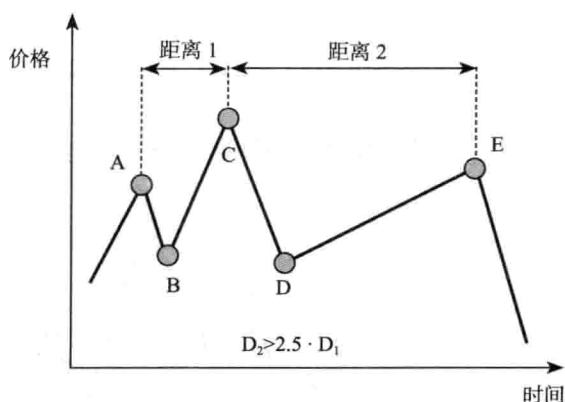


图 3-13 不符合标准的横向对称

完整形态法则：设定突破颈线的最大倍数

C&O 还提出了这样一个法则，即一旦右肩，也就是图 3-15 中的 E 点形成以后，向颈线移动的时间过长，以至于无法最终突破颈线的图形是

不包含在该法则当中的。在其他条件不变的情况下，这一标准是通过图形的内部比例，而不是以其他固定的倍数单位的数值决定的。这么做的结果就使得该法则可以应用于所有形态之中，而不用考虑其倍数或者范围的因素。

最大倍数的概念考虑到从右肩（即 E 点）开始移动，直到突破颈线的距离应当与左右两个肩部（A 点和 E 点）到头部的距离相等。图 3-14 中的形态符合了完整形态法则中的倍数标准，因为从右肩到突破颈线（D4）的大致距离小于左右两个肩部（A 点和 E 点）之间的距离（D3）。然而图 3-15 中的形态就不符合这一标准，因为 D4 的距离大于 D3。

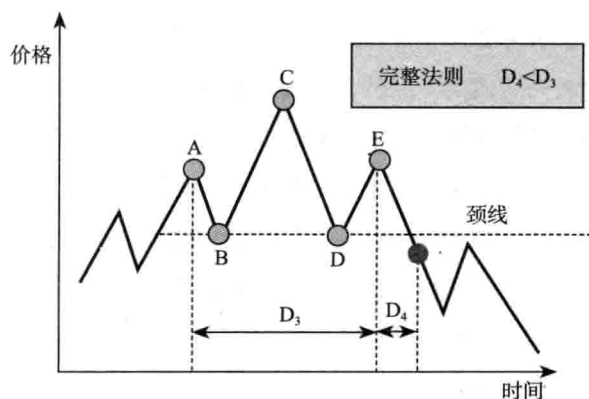


图 3-14 符合标准的完整形态法则

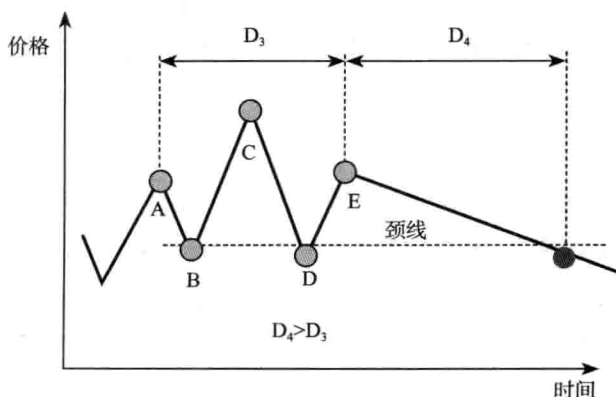


图 3-15 不符合标准的完整形态法则

未来信息的泄露：前视偏差

C&O 在进行头肩形态模拟的时候，对将来有可能会出现的信息泄露或者前视偏差问题做出了预防。他们假设当一项交易决策制定完成时，其中所包含的信息内容并不是可以获取的，这一问题一直困扰着回溯测试。因此，在进行回溯测试时，就有可能出现在回溯测试中得出的盈利值比实际交易所获得的盈利值要高得多的现象。一个最为极端的例子就是假设我们提前拿到了将要在第 2 天出版的《华尔街日报》。

在对头肩形态进行回溯测试时，如果 Z 字形的临界值百分比大于右肩和颈线之间的距离百分比，那么就有可能发生未来信息的泄露。但是我们也不能以此就假设卖空信号的出现就是由于在右肩（E 点）被发现之前已经发生的颈线突破所造成的。然而，在右肩通过“Z 字形过滤器”被确定之前，价格走势图很有可能已经与颈线交叉了。为了证明，我们假设右肩（E 点）所在的位置是位于颈线之上 4% 的地方，但是“Z 字形过滤器”的位置则位于颈线之上 10% 的地方。由于从右肩开始 10% 的下跌已经可以明确确认右肩的形成了，所以对于一个交易者知道颈线被突破的假设就显得不那么合理了。在下面的这个例子中，假设入市价格在右肩以下仅 4% 的地方，要比入市价格在右肩以下 10% 的地方更为合理，因为价格的走势图可以形成一个完整形态的头肩形态。为了避免类似的问题，C&O 假设入市价格为通过“Z 字形过滤器”对右肩进行客观的定义之后。尽管对入市价格的定义在某些实例当中的运用并不恰当，但是回溯测试并不会受到前视偏差的影响。我个人之所以注意到了这一点，是因为 C&O 对细节更为重视。

C&O 的形态定义法则还涉及其他的因素，例如，入市与离场位置、止损线等，但是关键的一点是它使得将主观的图形形态转变为客观的、可以检验的形态成为了可能。读者可以对 C&O 提出的任何一种形态提出质疑。希望读者可以发现更好的关于图形形态的客观定义。

最后，C&O 把他们自己用于自动确认头肩形态的运算法则介绍给了其他图形分析师。C&O 声称图形分析师们会认同客观的头肩形态计算法则实际上是符合主观的头肩形态确认标准的。

头肩形态回溯测试的结果

对于股票和货币来说，头肩形态是否真的具有可预测的信息吗？用一句话来概括，被称为图表分析的基石的这种形态本身就是错误的。由C&O进行的测试显示，头肩形态对于股票和货币市场赚取适度的盈利来说是没有意义的。头肩形态在对货币市场进行的6个测试中，只有其中的2个产生了盈利，而相关的更为复杂的头肩形态运算法则被更加简单的基于“Z字形过滤器”的客观信号所取代。此外，当C&O运用“Z字形法则”确定或者否定头肩形态出现的信号时，头肩形态并没有产生更多的信息价值。换句话说，当头肩形态信号出现在相同走势方向的时候，“Z字形信号”的表现并不好，而在头肩形态信号出现在相反走势方向的时候，“Z字形信号”的表现还是不错的。对于货币交易者来说，他们的底线就是：头肩形态的价值是难以预测的。⁴⁷

头肩形态在股票中的表现是较差的。C&O从1962年7月到1993年12月这段时间里随机挑选了100只股票，⁴⁸并对头肩形态进行了评估。从平均值来看，每一只股票每年会出现一次头肩信号，其中包括了多头和空头信号，总计超过3100个样本。为了检验头肩形态在实际股票价格中的盈利情况，C&O对于头肩形态在假设的历史价格中的业绩表现的基础之上建立了一个基准。这些模拟的历史价格是通过随机的方式把从实际的价格变化当中产生的历史价格汇集在一起的。运用实际的价格变化，假设的历史价格对于实际的股票来说，具有同样的统计学特点，但是任何对形态结构的技术分析更倾向于运用这种真实而又短暂的结构——尽管事实告诉我们假设的历史价格是随机产生的，而且头肩形态也确实符合C&O对于此种形态的定义，但是这种形态的预测力还是会被淘汰的。这一点验证了哈利·罗伯茨早期得出的结论。

如果出现在现实股票数据中的头肩形态是切实可用的话，那么它所产生的盈利就应当大于那些出现在假设的股票历史价格中的头肩形态所产生的盈利。C&O发现在实际股票价格中出现的头肩形态比出现在假设的股票历史价格中的头肩形态所产生的亏损要少一些。根据“结果一致显示运用头肩形态进行的交易是不具有盈利性的”的研究，头肩形态产生的信号在

持有期为 10 天的交易中, 平均亏损 0.25 个百分点。而这一数值在假设的股票数据中出现的头肩形态产生的信号中为 0.03%。C&O 把运用头肩形态进行交易的人称为“噪声交易者”, 即把一个随机的信号误认为是包含丰富信息的投机者。

对 C&O 的发现予以确认的是罗闻全 (A.W.Lo) 等人。⁴⁹ 罗在核回归 (kernel regression) 的基础之上运用了一种将头肩形态客观化的替代方法, 即一种复杂的局部平滑技术。⁵⁰ 他们的研究并不能够证明头肩形态是无用的⁵¹ 这一原假设。

布尔科夫斯基 (Bulkowski)⁵² 发现头肩形态是可以产生盈利的, 但是他的研究并不深入。他没有提出关于回溯测试的形态或者入市 / 离场法则的客观的形态定义。换句话说, 他的研究是主观主义的技术分析。另外, 研究结果没有把他对形态进行检验这一时期的股票市场的趋势进行调整。

技术分析的子集

基于前面的讨论, 技术分析包含 4 个子集: ①主观的技术分析; ②带有未知的统计学意义的客观的技术分析; ③没有统计学意义的客观的技术分析; ④含有统计学意义的客观的技术分析。

主观的技术分析, 已经被定义为不能得出回溯测试的方法。余下的 3 个子集均指客观的技术分析方法。

第 2 个子集包括未知数值的客观方法。尽管这些方法是客观的, 也是可以进行回溯测试的, 但是测试结果并没有对其统计学的意义做出评估。因此, 它将作为一种简单的应用方法在第 4 ~ 6 章进行讨论。这并不意味着在回溯测试的结论当中运用统计学的方法是简单的事情, 反而是做出决策所必须要做的工作。

我所说的第 3 个子集是指无效的技术分析, 包括客观的技术分析法则, 这些法则的结论都是经过全面综合的回溯测试以及统计学的方法检验的, 但是它们既没有在独立存在的情况下, 也没有在与其他的方法相结合运用时体现出更多有用的价值。在所有可能性当中, 大多数客观的技术分析方

法都属于这个子集,因为金融市场固有的复杂性与随机性使得对市场进行预测的难度大大增加。实际上,在所有科学领域中,大多数提出的思想都是无法进行实践的。而重要的发现也是不多的,这是因为大部分科学领域的失败是不会有科学期刊上刊登的,所以这一点看起来也不是很明显。我想最重要的办法就是适当淘汰一些无效的方法。

第4个子集是有效的技术分析。它包括含有统计学意义的客观的技术分析,或者具有更好的经济学意义结果的技术分析方法。尽管有些技术分析法单独运用起来会十分有效,但是在面对金融市场的复杂性和随机性时,对于它与别的法则结合而成的另一种更加复杂的方法的运用,也许会像大多数的法则那样产生些许的价值。

实证技术分析(EBTA)指的是子集(3)和子集(4),即经过了回溯测试以及符合统计学分析的客观的技术分析。根据我们前面所讨论的内容,关于技术分析的分类如图3-16所示。

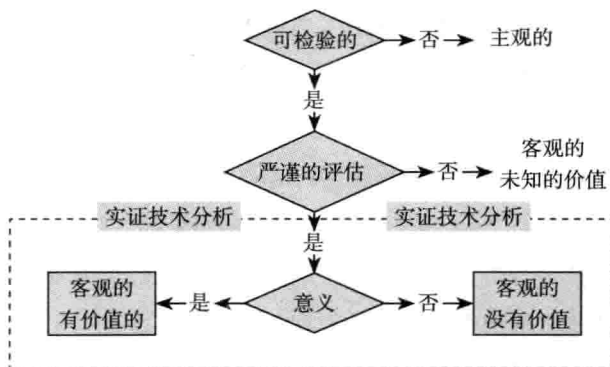


图 3-16 技术分析的子集

在接下来的3章中,我们将讨论统计学分析的应用以及回溯测试的结果。

第4章

Chapter 4

统 计 分 析

统计学是数据的科学。¹ 19 世纪末，享有盛誉的英国科学家、作家 H.G. 威尔斯（H.G.Wells, 1866—1946）曾经说过：20 世纪自由社会中睿智的公民必须理解统计学的方法。据此，我们也可以说 21 世纪技术分析的从业者和实践者也有同样的需求。本章及接下来的两章将讨论与技术分析有关的统计学内容。

当大量的数据能够清楚地表达出某种信息的时候，就不需要使用统计学的方法了。如果所有的人在饮用了一口水井里的水之后死于霍乱，而其他的人在饮用了不同水井里的水之后依然健康如初，那么我们会很明确地知道到底哪口井的水被病毒感染了，所以，我们此时同样不需要运用统计学的方法。然而，当数据所反映的内容存在不确定性时，统计学的方法就是最好的，也许是仅有的能够推断出合理的结论的方法。

确认哪种技术分析方法具有真正的预测力存在高度的不确定性。即使是最有效的技术分析法则在不同的数据集之中的表现也是不尽相同的。因此，统计学的分析是将有效的分析方法从无效的分析方法中分辨出来的唯一实用的方法。

不论技术分析的从业者是否承认这种法则，技术分析的本质就是统计推断。技术分析试图通过形态、法则等形式对历史数据进行归纳概括，以此对未来做出推断。外推法本身具有不确定性。这种不确定性使人感到担忧。

这种不安可以通过两个方面来解决。一方面，无视这种不安的存在。另一方面，就是通过统计学的方法，即承认这种不确定性的存在，并将其进行量化，然后再做出最好的决策来正视这种不确定性。英国知名的数

学家、哲学家伯特兰·罗素 (Bertrand Russell) 曾经说过,“不确定性,在鲜明的希望与恐惧面前,是会使人痛苦的,可是如果在没有令人慰藉的神话故事的支持下,我们仍希望活下去的话,那么我们就必须忍受这种不确定性”。²

很多人对统计学的分析持怀疑,甚至是鄙视的态度,而统计学家们通常把自己描述成为与现实生活脱节的只会乏味地捣鼓数字的傻瓜。对于这种说法我们把它当作笑话来讲吧。我们嘲弄自己并不理解的东西。有一个故事讲述的是一个身高 6 英尺的男人淹死在平均深度只有 2 英尺的池塘里的事。还有一个则关于 3 个统计学家去猎鸭子的故事。有 3 个统计学家朝着从头顶飞过的鸭子开枪,第一个统计学家开枪了,但是子弹偏左了大概 1 英尺。第二个统计学家也跟着开枪了,同样没击中,子弹偏右了 1 英尺。第三个统计学家跳起来兴奋地嚷道:“嗨,平均来讲,我们打中了!”即使是平均的错误率是零的话,那这三个人也不会在晚餐中见到鸭子了。

强大而有效的工具有可能被用于不怀好意的目的。批评家们常常指责统计学被用于歪曲事实和欺骗的目的。当然,相似的结果也有可能用言语来表达,尽管语言并不是那么好掌握。我们需要更为理性的态度来面对这一切。与其用怀疑或者只看表面现象的态度考虑所有基于统计学得出的观点,“倒不如用更为成熟的态度最大限度地学习统计学,并且将真实的、有用的结论从欺骗和愚蠢当中分辨出来”。³ 不加区分就接受统计学的人容易上本不该上的当。但是,因为不谨慎而拒绝统计学的人经常会显得无知。处于盲目拒绝和盲目轻信中间的人是那些思想开放的怀疑主义者。我们需要能够熟练解释数据的能力。⁴

统计学推理概览

对于技术分析的从业者和实践者来说,统计学推理是一个全新的领域。只要你知道去一个陌生的地方旅行的目的,旅行就会变得更加轻松。我们现在对接下来 3 章的内容进行预览。

通过在第 3 章中进行的论证,我认为对所有的技术分析都不具有预测

力，并且通过回溯测试得出的盈利完全是由于运气的原因进行假设来开始该章的内容是非常明智的。我们把这种假设称为“原假设”。运气，在当前的情况下，意味着在经过对某种法则进行测试之后的历史数据样本中，对该法则发出的信号和之后的市场趋势之间的反应既是顺理成章的，也是偶然的。尽管这个假说可以作为合理的出发点，但是也为运用实证对此进行反驳打开了一扇门。换句话说，如果观察资料与原假设所做出的预测相悖的话，那么它就会放弃这种假设，并且接受可以产生预测力的备择假设。在对法则进行测试的环境下，对原假设进行反驳的实证证据就是经过回溯测试检验的收益率，因为这个收益率之高，已经无法合理地将其归结于仅仅是运气的原因了。

如果技术分析法则不具有预测力，那么其在消除趋势⁵的数据中的预期收益率将为零。然而，在任何一个小型的数据样本中，没有预测力的法则的盈利能力会都严重地与数值0相背离。这些背离的数值就是运气的最直接的表现，即好运气与坏运气。这一现象我们可以在掷硬币的实验中得以重现。在为数不多的掷硬币过程中，印有头像的那一面出现的概率可能会严重偏离0.5这个数值，但是0.5这个值在投掷次数很多的情况下，是人们期待出现的。

一般来说，无效的法则与零收益率之间背离的机会是很小的。然而，有时候，一种无效的法则单凭好运气就可以产生巨额的收益。这些极为罕见的例外会误导我们相信一种无效的法则是具有预测力的。

避免受到这种误导影响的最佳办法就是了解这种靠运气获得的收益的程度如何。这个任务需要由数学函数来协助完成，而这个数学函数的工作就是说明在偶然的情况下出现的收益与零收益之间的离差值。这些就是统计学能够为我们做的工作。

我们把这种函数称为“概率密度函数”，它给出了每一个与数值0相背离的正/负偏差的概率。换言之，它显示了在哪种情况下可以导致无效的法则产生盈利的比率。图4-1显示了一个概率密度函数。⁶图中密度曲线都集中在数值0的现象反映了原假设坚持认为该法则的预期收益率为零。

图4-2显示了经过回溯测试获得的正收益率。这为我们提出了这样一

个问题：我们观察到的正收益率是否足以保证拒绝接受某种法则真实的收益率为零的原假设呢？如果被观察到的绩效正好落入仅仅是由于偶然而形成的离差值范围之内，那么实证就无法有力地反驳原假设了。在这种情况下，原假设要经受回溯测试证据的实证挑战，因此，我们根据证据所做出的谨慎解释就是该法则并不具有预测力。

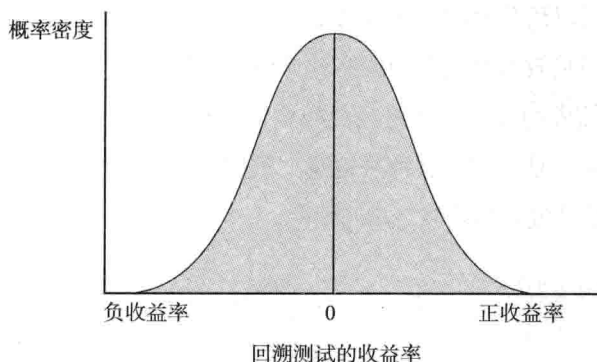


图 4-1 偶然情况下的绩效的概率密度，即无效的技术分析法则的大概绩效

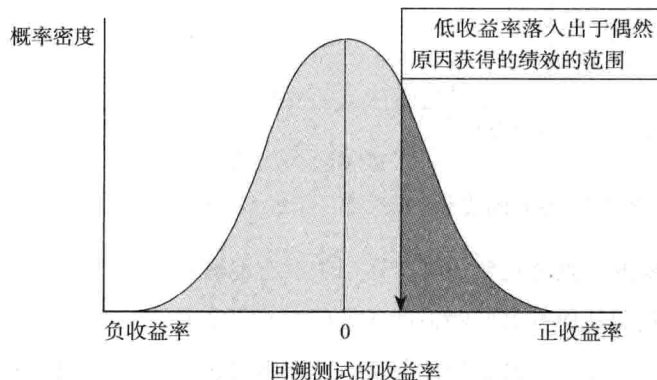


图 4-2 无效的技术分析法则在偶然情况下获得绩效的概率

回溯测试证据的取值范围由概率密度函数的部分区域⁷进行量化，其数值可以等于或者大于该法则已被观察到的绩效的值。在图 4-2 中，密度函数的比例由垂直箭头右边的阴影部分表示。这部分的规模可以理解为在偶然的条件下，由没有预测力的法则（预期收益率 = 0，或者原假设是真实的）产生较高的收益率的概率。当阴影占据了密度曲线相对较大的部分时，则意味着出于偶然获得正数绩效的概率也会得到同等的加大。在这种情况下

下，我们就不能对由原假设得出的结论做出合理的解释了。换句话说，就是无法对没有预测力的法则所得出的结论进行辩解了。

然而，如果被观测的绩效距数值 0 的距离很远的话，那么概率密度函数所占极端数值的比例就会很小，则产生正收益的绩效就与声称该法则不具有预测力的观点相矛盾。换句话说，就是这一证据可以对原假设进行反驳。我们还可以用其他的方法来考虑这个问题：如果原假设是真实的，则产生正收益的绩效水平出现的概率就应该非常低。这个概率是由等于或者大于被观测绩效的数值位于密度曲线中的比例来表示的。我们通过图 4-3 来显示。这并不是说如果这种法则缺少预测力，那么位于密度曲线最右边尾部的被观测绩效将会是与之相匹配的。

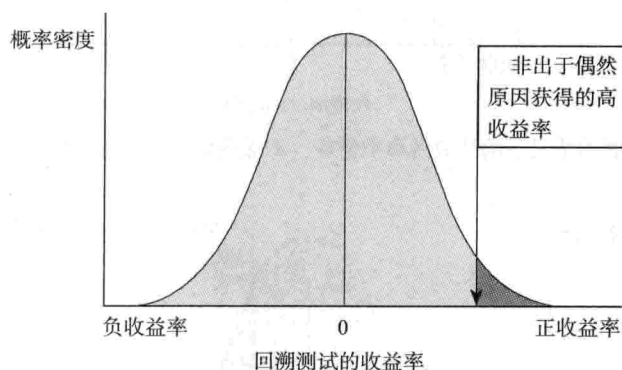


图 4-3 有效的技术分析法则在偶然情况下获得绩效的概率

最重要的一点就是要知晓实证当中没有能够传递给我们的内容。实证没有告诉我们任何有关原假设和备择假设是真实有效的概率。它只是说明了在原假设实际上是真实的情况下，实证将会出现的概率。因此，实证将会出现的可能性，并不能代表这个假设就是真实的概率。在原假设是真实有效的情况下，被观察的实证的出现是不大可能的，我们也可以由此而推论出这个原假设就是错误的。

现在回顾一下第 3 章中曾经提到的“有四条腿的动物就是狗”这句话，这个说法是无法确定这种假设的真实性的。尽管“有四条腿的”这个证据与它就是狗的假设存在一致性，但是从演绎推理法的角度来说，这并不足以证明它就是狗。同理，当被观察的实证的正收益绩效与假设所说的该法

则具有预测力的说法相一致时，也是不能证明这种法则就是具有预测力的。目前存在着一种争论，即试图运用实证来证明某种假设的真实性，而这些实证与承认确认结果的逻辑错误的假设是一致的。

如果这种动物是狗，那么它有四条腿。

这种动物有四条腿。

无效的结果：因此，这种动物就是狗。

如果一种法则具有预测力，那么对它进行回溯测试就会产生盈利的结果。

回溯测试是会产生盈利的。

无效的结果：因此，这个法则具有预测力。

然而，如果缺少“它有四条腿”这个信息的话，我们就很容易证明“这种动物是狗”这个假设是错误的。⁸换句话说，被观察的实证可以用来证明一种假说是错误的。这种论据会应用于有效的演绎推理形式当中，即对结论的否定。这种论据的一般形式是对结论的否定，即：

如果 P 是真实的，那么 Q 也是真实的。

Q 不是真实的。

有效的结果：因此，P 是不真实的（例如，P 是错误的）。

如果这种动物是狗，那么它就有四条腿。

这种动物没有四条腿。

有效的结果：因此，这种动物是狗的说法就是错误的。

上面所给出的“没有四条腿”这个论证，最终证明了这种动物是狗的结论是错误的。然而，这种程度的确认并不能为科学和统计学所接受。有些论证也许永远无法证明某种假设是错误的。不过，如果一种假设是真实的，那么类似的逻辑可能会表明这些实证是难以相信的。换句话说，实证赋予了我们质疑假设的权利。因此，具有高盈利特征的回溯测试可以向认为该法则没有预测力（例如，预期收益率为零）的假说提出质疑。

如果一种法则的预期收益率等于零或小于零，那么回溯测试

所产生的盈利就会接近零。

回溯测试的绩效并没有接近零；实际上，它的值超过了零。

有效的结果：因此，这种法则的预期收益率等于零或小于零的论点很有可能就是错误的。

用回溯测试结果为准的绩效这一事实来证明一种法则缺少预测力是件让人多么无法相信的事情啊。目前有力而又快捷的法则是不存在的。按照惯例，在原假设是真实的情况下，除非被观察的绩效出现的概率在 0.05 或者更少，否则大多数的科学家是不会否定这种假设的。这一估值被称为观察的统计显著性。

到目前为止关于案例的讨论仅限于对单一法则进行回溯测试。然而，在实际中，对技术分析法则的研究绝不仅仅局限于单一法则。经济学的计算能力、多功能的回溯测试软件以及大量的历史数据使得这种研究变得容易，我们把这种以选择一个最佳绩效为目标，对多种法则进行测试的实践称为数据挖掘。

尽管数据挖掘是一种非常有效的研究方法，但是对多种法则进行的测试会使得出的绩效增加运气的成分。因此，用于否定原假设的绩效的临界值就必须设定得更高一些，或许还要更高。这些设定的较高的临界值是为了弥补在回溯测试中由于幸运成分所导致的无效法则出现的可能性。有关数据挖掘偏好的话题，我们将在第 6 章进行讨论。

图 4-4 比较了两种密度函数的概率。上面的图适合对单一的法则进行回溯测试，并且计算其显著性。而下面的图则适合对进行回溯测试的 1 000 种法则中绩效最佳的那一种法则进行显著性的计算。下图当中的密度曲线将从数据挖掘当中不断增加的幸运成分的概率考虑了进来。有一点需要注意的是，如果对这种最佳法则进行观察的绩效是通过仅仅适用于单一法则的密度曲线计算得出的，那么由于其绩效已经远离分布图中最右边的尾部，因此就会产生显著性。然而，如果对这种最佳法则进行观察的绩效是通过其概率范围适当的密度函数计算得出的，那么统计上的显著性就不会显现。也就是说，这种法则的绩效值越高，就越不能够保证这种法则具有预测力，或者其预期收益率大于零的结论。

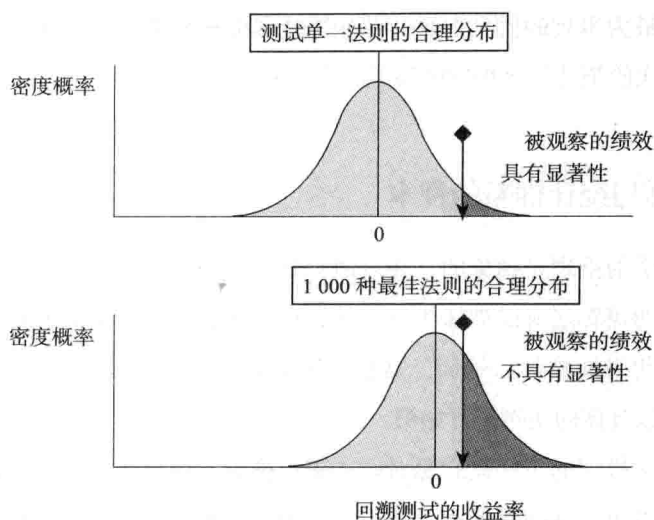


图 4-4 对单一法则进行回溯测试所产生的优异绩效，
在 1 000 种被测试的法则当中仅仅属于中等水平

严谨的统计学分析的必要性

一门学科所采用的工具和方法，对其能够发现的东西是有局限性的。将它们进行改进之后就可以获得更多的知识。随着望远镜的发明，天文学取得了质的飞跃。尽管那时的望远镜与今天的望远镜相比确实粗糙了很多，但是这些早期工具的分辨率却是肉眼的 10 倍。技术分析也具有类似的机会，但是它必须用严谨的统计学分析代替那些非正式的数据分析。

非正式的数据分析是无法完成从金融市场中获取有效信息的任务的。虚幻的图形形态中的数据被人们所重视，而真实的图形形态却因为噪声和复杂性被忽视。严谨的统计学分析对于这个困难的任务来说是再合适不过了。

统计学的分析是一套定义完整的，对数据进行收集、分析以及解释的程序。本章以及接下来的两章内容将会介绍统计学工具的工作和推理方法，并以此来确定有效的技术分析法则。我们在这里的概述会很简要，在很多实例中，为了便于理解，我把数学的严谨性暂且放下。然而，这并不意味

着削弱了最为本质的信息内容，即如果技术分析想要自圆其说的话，它就必须用正式的统计学分析这种科学的方法来证明。

抽样与统计推断的样本

统计学的推理是抽象的，并且通常是与常识相违背的。与非正式推理背道而驰的逻辑之所以被认为是好的，是因为它可以帮助我们把这些根深蒂固的思想彻底拿掉。然而，这也就是它让人难于理解的地方了。所以，我们还是以具体的实例来开始吧。

统计学推理的主要概念是对样本进行推理。对某个观察资料的样本进行研究，得出一种图形形态，然后这个图形形态就会被赋予与被观察的样本之外的某个事件相适应（由此推断出来的）的预期。例如，在对历史样本中的法则进行盈利性的观察时，这种法则就被赋予了未来盈利的预期。

我们现在来考虑一个与技术分析毫无关联的问题。这个问题来自瓦利斯（Wallis）和罗伯茨（Roberts）⁹ 所合著的《统计学，一种新的方法》（*Statistics, A New Approach*）一书。这个问题所关注的是一个灰色和白色珠子混合在一起的盒子。在这个盒子中，全部珠子的数量以及灰色和白色珠子的数量都是未知的。而我们的任务就是确定盒子当中灰色珠子的数量。为了简便起见，我们把灰色珠子的数量用 F-G（fraction-grey in box）表示。

我们用一种便捷的方式使这种情况与实际当中所面对的统计学问题类似。我们不允许在同一时间内对盒子的内容进行观察，这也就避免了对灰色珠子的数量（F-G）进行直接的观察。这一限制使得这个问题变得愈发的真实，因为在实际情况下，对诸如盒子中所有珠子等所有感兴趣的项目进行观察是不可能的，也是不实际的。其实，这种限制催生了对统计学推理的需求。

尽管我们不允许对盒子内容做全面的分析，但是我们可以每次抽取其中的 20 个珠子并对它们进行观察。所以我们对于获得灰色珠子的数量（F-G）的策略就相应地转变为在大量的样本中对灰色珠子的数量（F-G）进行观察。对此，我们共挑选 50 个样本进行观察。我们用小写的 f-g 代表实

验当中灰色珠子的数量。

样本是通过以下的步骤获得的：我们在盒子的底部安装了一个包含 20 个小坑的可以滑动的托盘，每个小坑的大小与盒子当中的珠子的大小相匹配。托盘可以在手的推动之下滑出盒子，而滑出部分则由相同的托盘所替换，这就使得剩下的珠子不会因为托盘的滑动而掉到盒子外面。因此，当托盘每一次滑动的时候，我们就能够获得装有 20 个珠子的样本，并且有机会对 $f-g$ （实验当中灰色珠子的数量）进行观察（见图 4-5）。

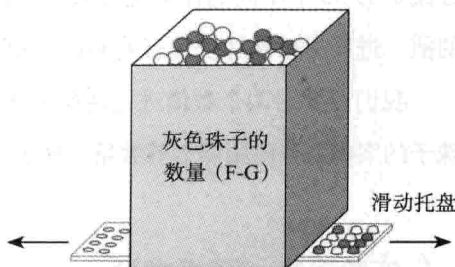


图 4-5 确定每一个样本中 $f-g$ （实验当中灰色珠子的数量）

在确定所给的样本中 $f-g$ （实验当中灰色珠子的数量）之后，我们把这 20 个珠子放回盒子中并重新摇晃后，再进行下一组实验。这样就使得盒子当中的每一个珠子都有同等的机会在下次实验中被摇出。用统计学的说法，即我们确定每一次的实验都是随机的。我们将每一次实验所获得的 $f-g$ （实验当中灰色珠子的数量）记录下来，然后把珠子放回盒子中并重新摇晃后，再进行下一组实验，我们将这个过程重复 50 次。最后，我们会得到 50 个不同的 $f-g$ （实验当中灰色珠子的数量），每一个数值代表每一组样本。

在将每组得出的样本放回盒子并进行下一组实验之前，我们会得到一个稳定的 $f-g$ （实验当中灰色珠子的数量）集中度。也就是说，灰色珠子的数量 $(F-G)$ 在这 50 次实验当中始终会保持一个常量。随着时间的推移，统计学依然保持着稳定的特性，我们就会将其认为是稳定的、固定的。由于许多组珠子会完全地从盒子当中取出，如果在每组样本之后没有将这些珠子进行替换，那么灰色珠子的数量 $(F-G)$ 就会随着 50 次取样的推移而改变。如果统计学的特性随着时间的推移而改变，我们就将其称为非稳定的。

记住灰色珠子的数量 $(F-G)$ 与 $f-g$ （实验当中灰色珠子的数量）之间的区别是非常重要的。 $F-G$ 是指在整个盒子当中的灰色珠子的数量。用统计

学的语言来说,即我们把所有感兴趣的观察资料称为**总体**。在这个样本当中,术语“**总体**”是指在盒子当中所有颜色的珠子。术语**样本**是指“**总体**”的子集。因此,F-G是指“**总体**”,而f-g是指“**样本**”。然而,我们的任务就是对50个不同的样本进行观察,得出f-g(实验当中灰色珠子的数量)的值,进而尽可能地获得有关F-G(灰色珠子的数量)的数值。

我们还要对两个数值进行明确的区分,即包含样本的观测数量(20个珠子的案例)与使用样本的数量(50次实验的案例)。

实验概率与随机变量

概率,就是指有一定几率的数学。¹⁰ **实验概率**是对没有明确结果的环境做出观察或者处理。¹¹ 它包括计算长途旅行后返程到达的精确时间、观察暴风雪中降雪厚度或者观察在掷硬币过程中出现正面(头像)的行为。

我们把在实验概率中对数量和质量的观察称为**随机变量**,它包括掷出硬币后所显示出的正面或者反面、降雪厚度以及对观察资料的样本做出概括(例如,样本均值)。我们之所以把被观察的数量和质量称为是随机的,是因为它受到了偶然因素的影响。然而,对随机变量进行独立的观察是不具有预测性的,但是根据定义,对随机变量进行大量的观察是可以使其具有高度的预测性特征的。例如,在掷硬币的过程中,我们无法预测硬币掉下来之后会出现正面还是背面。但是,我们将这个过程重复1 000次,那么出现正面(头像)的次数就会确定在500次这个范围内,这个数值可以说是具有高度的预测性了。

根据定义,一个随机变量可以对两个不同的数值进行假设,尽管可以假设的更多——也可以是无限的数字。存在于掷硬币过程中的随机变量,在掷完之后出现的任何一面,可以假设为两种可能的值(正面或者背面)。对于在中午时分自由女神像基座下面温度的定义,我们可以把它的值假设得更大一些,也可以设定在温度计的精确数字之内。

抽样：最重要的实验概率

在统计学当中最重要的实验概率就是抽样。它包括从总体当中提取出观察资料的子集。我们把目前存在争议的随机变量称为样本统计量。样本具有可以计算的特点。 $f-g$ （实验当中灰色珠子的数量）就是样本统计量的最好范例。

我们可以通过很多种方法来进行抽样，但是在抽样中重要的一点就是对样本的选择必须是通过随机的方式进行的，而且样本间都是独立存在的。这就意味着在样本中的所有观察资料都会有一个平等的机会来完成此次实验。由于样本当中的观察资料之间的关系都是平等的，所以能够完成实验的观察资料都是出于偶然的因素。样本的创立必须以这样的方式进行是因为概率的原理，而这一原理又是依赖于资料统计学推理的，它是以观察资料会以随机的形式完成此次实验为基础的。

我们来分析一下对志愿者进行新药测试的样本。这个样本的构建必须是在让接受药品实验的任何一名志愿者都具有被选入测试小组的平等机会的基础之上。如果接受实验的人不是通过这种方式挑选的，那么测试的结果就会导致得出不正确的结论。在药品实验中的那些被挑选的志愿者，假设他们被赋予对药物测试做出良好反应的预期，那么这个实验就会导致结论偏差。换句话说，他们对这种药物对于一般大众的疗效将会非常的乐观。

试想一下，你闭上眼睛从盒子里随意拿出一颗珠子，然后观察它的颜色：这个抽样实验所采用的是单一样本。珠子的颜色是灰色或者是白色，是随机变量。现在，试想一下从盒子里随机挑选 20 颗珠子。这个抽样实验的样本规模是 20，而 $f-g$ （实验当中灰色珠子的数量）的样本统计量仍然是随机变量。

正如我们之前讨论的，我们把这种类型的随机变量称为 $f-g$ （实验当中灰色珠子的数量）。这个随机变量是受到偶然的因素影响的。我们假设 21 个不同的数值（0,0.05,0.10,0.15,...,1.0）。现在我们回到主题：增加有关 $F-G$ （灰色珠子的数量）的数值的知识。

从样本中推测出的知识

假设第一个样本中包含 20 颗珠子当中有 13 颗灰色珠子, 那么这个样本统计量的 $f-g$ (实验当中灰色珠子的数量) 数值就是 0.65 ($13/20$)。我们从这个数值当中能够得出多少有关 $F-G$ (灰色珠子的数量) 的内容呢? 仅靠这点信息, 也许有人就会试图得出结论, 认为我们已经解决了这个问题, 即 0.65 这个值就是盒子中所有灰色珠子的数量 ($F-G=0.65$)。这种观点很自然地假设一个单独的样本就可以提供完全的知识。而与这种观点完全相反的是那些错误地认为该样本太小, 以至于不能够提供任何的信息的人。

以上的两种结论都是错误的。首先, 我们根本就没有从单一的样本中获得完整的知识的思想基础。即, 使包含有 20 颗珠子的单一样本有可能对整个盒子中的珠子进行完整的复制, 但这是不可能的。例如, 如果 $F-G$ (灰色珠子的数量) 的真实数据是 0.568, 那么一个包含有 20 颗珠子的单一样本是不可能得出这个数值的。 $f-g$ (实验当中灰色珠子的数量) 的值是 0.55 (其中有 11 个灰色的珠子), 也许这是与 0.568 这个值最为接近的数值了。样本中包括 10 颗珠子的 $f-g$ 值是 0.5, 包含 11 颗珠子的 $f-g$ 值是 0.55, 包含 12 颗珠子的 $f-g$ 值是 0.60。

而那些声称从样本中什么都不可能获得的人的想法也是错误的。单以这个样本为例, 上面的两种概率都可以被排除。我们可以排除 $F-G=1.0$ (所有的珠子都是灰色的) 的确认值, 因为这里面包括 7 颗白色的珠子。同理, $F-G=0$ 的主张也是错误的, 因为样本里面包含着 13 颗灰色的珠子。

然而, 以单一被观察的 $f-g$ (实验当中灰色珠子的数量) 的值为基础, 对 $F-G$ (灰色珠子的数量) 的精确数量做出的估计会趋于不确定性。即使我们采用大量的样本并得出了大量的 $f-g$ (实验当中灰色珠子的数量) 的值, 那么 $F-G$ (灰色珠子的数量) 的值仍然会存在不确定性。这是因为根据定义来说, 即一个样本用部分内容代表了整个盒子所包含的内容。只有对这个总体当中的每一颗珠子进行观察, 才能够消除对 $F-G$ (灰色珠子的数量) 的值的的所有不确定性。但是这种做法是被禁止的。

但是, 根据我们已经从单一的样本中所得到的内容来看, 关于 $F-G$ (灰色珠子的数量) 的值既不可能是 0 也不可能是 1.0 的某些知识, 有可能被

一些睿智的猜测所超越。例如, 尽管对 F-G (灰色珠子的数量) 的值为 0.10 的主张并不能够被最终排除; 但是没有人会认真地去对待它。从第一个样本中观察到的数值为 0.65 的 $f-g$ (实验当中灰色珠子的数量), 因为与 0.10 的差距太大而被认为不能对 F-G (灰色珠子的数量) 的值做出精确的评估。如果 F-G (灰色珠子的数量) 的值真的像 0.10 这么低, 那么按照常理来讲, 取 $f-g$ 值为 0.65 的这个样本的可能性就不太可能了。同理, 我们可以运用同样的逻辑将 $F-G=0.95$ 的主张也排除掉。如果整个盒子里的灰色珠子所占的百分比为 95%, 那么 $f-g=0.65$ 的这个读数也未免太低了。从本质上说: 单一的样本是可以提供一些有关 F-G (灰色珠子的数量) 的知识的。

我们从样本 50 当中学到了什么

关于 F-G (灰色珠子的数量) 更多的知识可以通过对更多的样本进行分析而获得。假设其他的 49 个样本已经计算出 $f-g$ (实验当中灰色珠子的数量) 的值。我们首先要注意的就是在不可预测的情况下, 各个样本之间 $f-g$ 值的变化。这种随机行为的特殊形式在统计学中是最为重要的一种现象。我们把这种现象称为**抽样变异性** (sampling variability) 或者**抽样变异** (sampling variation)。

抽样变异作为统计学结论当中的不确定性是非常重要的。抽样变异性越大, 由此导致的不确定性就越大。每个样本中 $f-g$ (实验当中灰色珠子的数量) 的值的波动越大, 对于 F-G (灰色珠子的数量) 的不确定性也就越大。

令人遗憾的是, 抽样变异性这个重要的现象, 对于分析数据的人来说, 是比较陌生的。这一点是可以理解的, 因为现实世界当中的问题没有提供除独立样本之外的其他样本。而“盒子里面的珠子”这个实验则提供了这个机会。

在我们将要看到的样本 50 实验中, 随机变化在 $f-g$ (实验当中灰色珠子的数量) 当中是非常明显的。这与法则的研究者所面临的环境是不一样的。他们通常拥有单独的市场历史样本, 如果用这些数据对法则进行测试, 我们只能获得一个被观察到的绩效统计数字。如果将这个法则在大量的独

立样本中进行测试的话，我们就没有机会看到该法则的绩效是如何变化的。从本质上来讲，抽样变异对于那些不熟悉统计学分析的数据分析师来说是不明显的。这是一个很大的盲点！

抽样变异可以在 50 个 f-g 数值的表格中看到，如表 4-1 所示。

表 4-1 样本 50 当中灰色珠子的数量

样本编号	f-g	样本编号	f-g	样本编号	f-g
1	0.65	18	0.60	35	0.60
2	0.60	19	0.40	36	0.55
3	0.45	20	0.60	37	0.50
4	0.60	21	0.45	38	0.50
5	0.45	22	0.55	39	0.45
6	0.45	23	0.45	40	0.50
7	0.55	24	0.50	41	0.55
8	0.40	25	0.70	42	0.45
9	0.55	26	0.55	43	0.55
10	0.40	27	0.60	44	0.60
11	0.55	28	0.70	45	0.60
12	0.50	29	0.50	46	0.60
13	0.35	30	0.70	47	0.65
14	0.40	31	0.50	48	0.50
15	0.65	32	0.90	49	0.50
16	0.65	33	0.40	50	0.65
17	0.60	34	0.75		

不出意外，f-g 的数值在不同样本间的波动是个随机变量。这属于在操作当中的偶然因素的样本。然而，f-g 的数值并不是完全由偶然的因素所决定的。f-g 的数值也同样受到 F-G（灰色珠子的数量）的数值的影响。

因此，我们可以说对 f-g 每一次观察的数值都是受两方面影响的结果，即 F-G（灰色珠子的数量）的表面现象，以及由抽样导致的随机性。F-G 发挥着重力中心的作用，它使得 f-g 的数值始终处于一个固定的范围。在某些样本中，随机性使得 f-g 的数值高于 F-G 的数值。而在其他的样本中，随机性使得 f-g 的数值低于 F-G 的数值。f-g 的数值以这个重力中心为基础，在不同的样本间随机波动。我们将这种现象称为抽样变异性。

另外一个非常重要的方面就是观察资料的数量影响着 F-G 的数值的不

确定性的水平程度。包含抽样的观察资料的数量越多，那么通过 $f-g$ 的数值所反映出的 $F-G$ 的数值就越精确。我们假设用 200 颗珠子的样本替代 20 颗珠子的样本。那么以 $F-G$ 为中心，出现在 $f-g$ 当中的随机变量就会变小。实际上，这些数量上的变化只相当于在 20 颗珠子的样本中所出现情况的 $1/3$ 。这一点非常重要：样本的规模越大，受到随机性的影响就越小。200 颗珠子降低了任何一个单独的珠子把 $f-g$ （实验当中灰色珠子的数量）从 $F-G$ （灰色珠子的数量）当中推离的能力。在 1 颗珠子的实验中，被单独选择的珠子所得出的 $f-g$ 的数值要么是 0，要么是 1.0。在 2 颗珠子的实验中， $f-g$ 的数值为 0、0.50 或者是 1.0。在 3 颗珠子的实验中， $f-g$ 的数值为 0、0.33、0.66 或者是 1.0。因此，样本的规模越大，在 $f-g$ 当中随机变量的数值就越小。大样本提供 $F-G$ （灰色珠子的数量），我们希望知道的事实是其反映自身的能力。这就是**大数法则**的作用，即大样本降低了偶然性的作用。换句话说，它表明了样本的规模越大， $f-g$ 的数值对于 $F-G$ 的数值的集中度就越紧密。这一点在统计学当中是最为重要的原理之一。

我们把已经学习到的重要的统计学概念用以下的内容进行表述：在抽样实验 50 的过程中， $F-G$ 的数值没有发生改变，而 $f-g$ 的数值在不同的样本中则会发生巨大的改变。我们把这种现象称为**抽样变异性**。这表明在任何情况下，一个随机的观察资料的样本是可以用来形成更大的领域（总体）的结论的。抽样变异性是由统计学推理¹²所提出的不确定性的来源。

抽样统计数字 ($f-g$) 的频率分布

在“盒子中的珠子”这个实验中，当抽样不包含灰色的珠子时，随机变量 $f-g$ 可以在从 0 开始的范围内，被假设为 21 种不同的数值，而当抽样中包含的珠子全部是灰色的时候，这个数值为 1.0。我们把 50 个被观察的 $f-g$ 的数值列在表 4-1 当中。我们在一次偶然性的分析中发现，有些 $f-g$ 的数值的出现概率比其他数值出现的概率要大。在 0.40 ~ 0.65 这个范围内的数值出现频率较高，而低于 0.40 和高于 0.65 的数值则从来没有出现过。在我们所采用的 50 个样本中，数值 0.50 出现了 9 次，0.55 出现了 8 次，0.60 出现了 10 次。请注意，在第 1 个样本中出现的数值 0.65，仅仅出现了 5

次。因此,那些没有出现在最常见的范围内的数值,并不代表着其就是不同寻常的。

一种叫作频率分布或者频率柱状图的平面图对这个信息的表述要比语言或者图表的形式更具有说服力。它显示了在样本 50 实验中每一个 $f-g$ 数值的出现频率。术语分布的说法是适当的,因为它显示了在随机变量中的一组观察资料是如何在变量可能出现的数值范围内分布的。

沿着频率分布的水平线轴的排列是一组区间的数列,或者说是二进制数列,每一个数字代表着一个可能出现的 $f-g$ 数值。每一个区间的垂直柱体的顶点代表着 $f-g$ 数值出现的次数。图 4-6 显示了样本 50 实验中 $f-g$ 数值的频率分布。

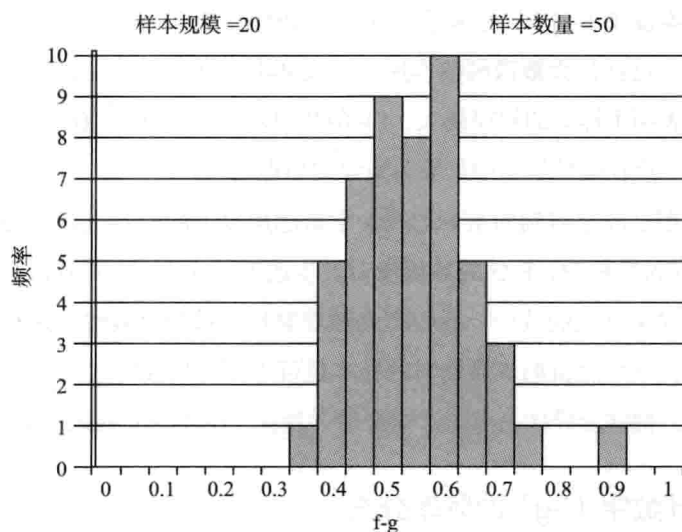


图 4-6 频率分布: ($f-g$)

频率和面积的均等

频率分布显示了特定的数值或者一组数值的频率和代表那些数值或者一组数值的柱状图所覆盖的区域之间的重要关系。如果我们对全部的分布进行分析的话,那么这个概念便可以得到理解。在图 4-6 当中,所有包含频率分布的灰色柱子代表了 50 个观察资料的 $f-g$ 值。因此,被所有垂直的柱子所覆盖的面积就代表了 50 个观察资料。你可以通过统计与所有垂直的

柱子有关系数量来证明这一点。而它们的统计总量也会是 50。

相同的原理可以应用于任何分布区域的分数。如果你要确定由单一的柱子所覆盖的频率分布的全部面积的分数，你将会发现柱子的分数面积与由柱子代表的全部观察资料的分数是相等的。例如，与 $f-g$ 数值 0.60 相关的垂直的柱子所显示的结果（频率）为 10。因此，这个柱子代表了被观察数值的 20%（ $10/50$ ）。如果想计算由这根柱子所覆盖的区域的价值，你就会发现它的面积占全部分布所覆盖的面积 0.20（见图 4-7）。

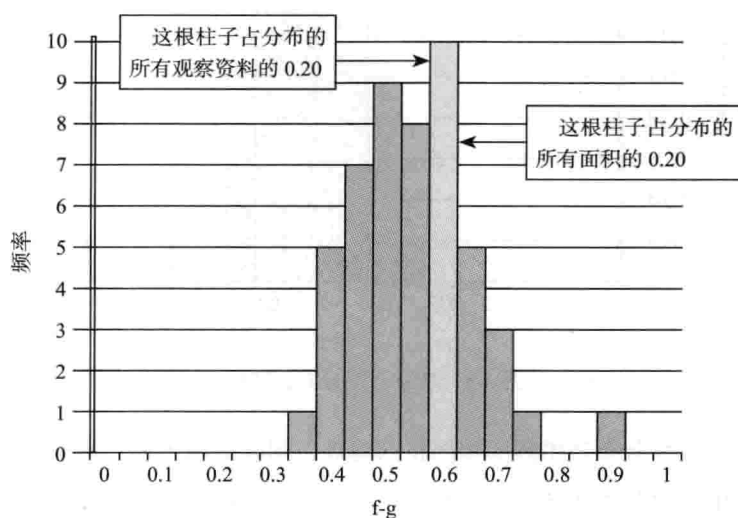


图 4-7 观察资料的比例与分布的全部面积相等

尽管这个观点看上去并无新意，甚至没有什么价值，但它是统计学推理的重要部分。最后，我们运用分布的面积变化分数来衡量在一种法则没有预测力的假设下，出于偶然的因素而在回溯测试中出现盈利的概率。当出现的概率很小时，我们就会得出这种法则具有预测力的结论。

$f-g$ 的相对频率分布

相对频率分布与之前讨论的普通频率分布相类似。在普通频率分布中的柱子的顶部代表了一个绝对的数字，或者说是被观察的随机变量的特殊值出现的次数。在相对频率分布中，柱子的顶端代表着被观察资料相对于（除以）包含分布在内的全部观察资料的数量而出现的次数。例如，假设 $f-g$

的数值为 0.60, 即在 50 个数据当中出现了 10 次。因此, 数值 0.60 就具有一个相对的频率 0.20, 即 (10/50)。数值 0.60 的柱状分布图就会向垂直范围的顶端值 0.20 靠近 (见图 4-8)。

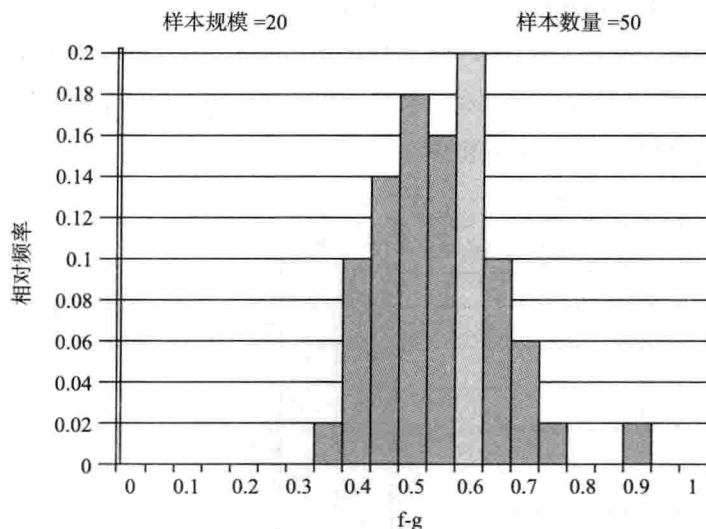


图 4-8 相对频率分布: (f-g)

普通频率分布当中频率和面积的均等, 同样适用于相对频率分布。所有柱子的相对频率等于 (总计) 1.0。简单地说, 即一个随机变量通常 (此时为 1.0) 会出现在由所有的柱状图所涵盖的范围之内。如果你把与所有包含分布的独立的柱子有关的相对频率相加的话, 那么结果就是 1.0 (100% 的观察资料)。因为每一个柱子代表着落入其区间的观察资料的部分, 根据定义来说, 即所有相加在一起的柱子所代表的就是所有的观察资料 (1.0)。

任何单一或者连续的柱子群的相对频率与它们在分布的全部面积中所占的比例是相等的。因此, 0.65 和更大的 f-g 值的相对分布就等于 $0.10+0.06+0.02+0.02=0.20$ 。这也就等于是说与 0.65 和更大的 f-g 值相关的连续柱子所重合的面积占分布总面积的 20%。当对技术法则进行测试的时候, 我们会做出类似的陈述。在此基础之上, 可以说等于或者大于 0.65 的 f-g 值的相对频率是 0.20 (见图 4-9)。

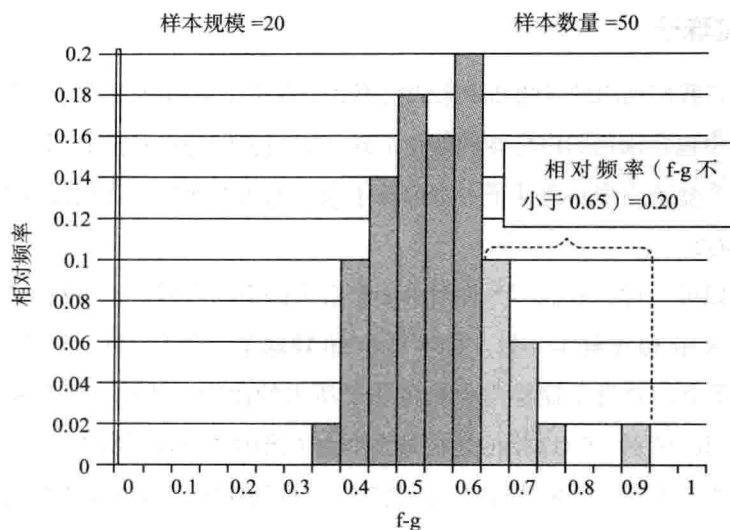


图 4-9 相对频率分布: (f-g)

我们从抽样当中获得了哪些关于 F-G 的知识

到目前为止, 我们通过抽样已经增加了有关 F-G 的知识水平。在第 1 个 20 颗珠子的实验中, $f-g$ 的值为 0.65。在此基础之上, 两种可能的结果被排除了, 即 $F-G=0$ 和 $F-G=1.0$ 。

除此之外, 我们采用了其他 49 个样本, 并对每一个样本的 $f-g$ 值进行观察。这 50 个样本表明 $f-g$ 的值在不同的样本间是随机变化的。然而, 除了某个不可预测的 $f-g$ 值, 在这些包含丰富信息的随机性当中产生了一个有组织的形态。这些分散的 $f-g$ 值通过合并组成了一个结构清晰的驼峰状形态。我们猜测这个集中的趋势与 F-G 的值有关, 而 F-G 的精确数值还是模糊的, 这是因为 $f-g$ 值当中的随机变量的原因。然而, 根据驼峰状形态的相对窄幅, 我们猜测 F-G 的值会位于 $0.40 \sim 0.65$ 这个范围之内。

基于我们开始的时候并没有关于 F-G 的值方面的知识, 进而无法对盒子中的内容进行全面的分析, 所以, 我们现在所获得的知识可以说已经相当多了。

第2盒珠子

现在我们利用抽样的知识来研究第2个盒子当中彩色珠子的比例问题。盒子当中包含灰色和白色两种颜色的珠子，而我们的目标是共同的，即获得第2个盒子当中灰色珠子占全部珠子数量的相对比例。我们把这个数量称为 $F-G2$ 。

和以前一样，我们不允许对第2个盒子的全部内容进行分析，但是我们可以采用50个样本，每个样本包含20颗珠子。 $f-g2$ 的值，指的是在第2个盒子里，对每个抽样当中灰色珠子所占的比例。然后，注意 $F-G2$ 与 $f-g2$ 之间的区别。 $F-G2$ 指的是在第2个盒子当中灰色珠子所占的比例，而 $f-g2$ 指的是在独立的抽样当中，灰色珠子所占的比例。 $F-G2$ 是不可观察的，而 $f-g2$ 则是可以的。

图4-10是 $f-g2$ 当中50个数值的相对频率分布。但是以下几点需要注意。

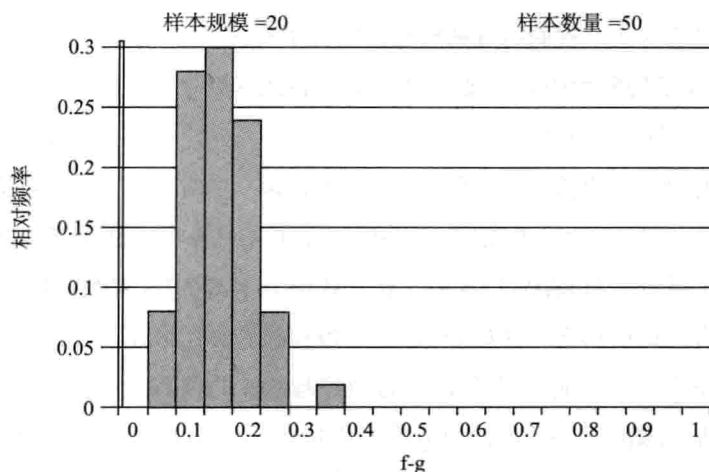


图4-10 相对频率分布，第2个盒子：($f-g$)

(1) 第2个盒子的一般分布情况与第1个盒子相似，其中以数值呈驼峰状聚集。

(2) 第2个盒子的中心数值的分布有别于第1个盒子。第1个盒子的中心值的分布靠近于0.55这个值，而第2个盒子的分布则接近于0.15这个

数值。我们在图 4-11 中可以清楚地看到。在图 4-11 中，我们看到了两个盒子的数值分布显示在了同一条水平线上。在它们每个中心值分布的上方，都用箭头做了标记。因此，这两个箭头代表了 $f-g$ 和 $f-g_2$ 的平均值。由此，我们可以推断第 1 个盒子比第 2 个盒子具有更高的灰色珠子的集中度。

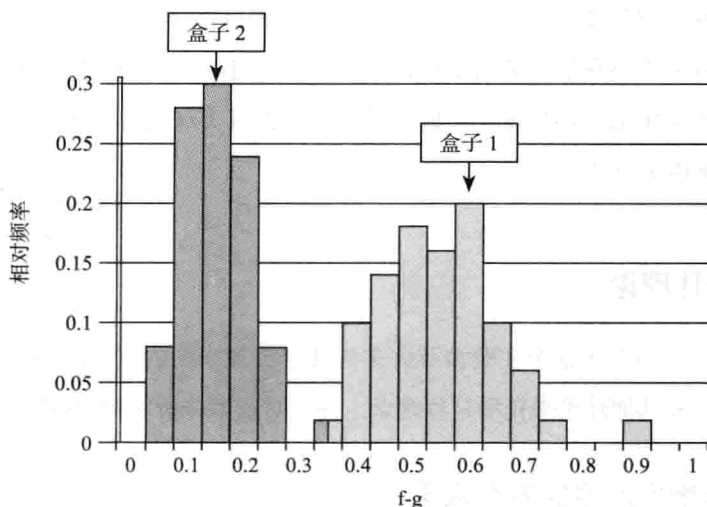


图 4-11 相对频率分布之比较

(3) 即使第 1 个盒子与第 2 个盒子当中的抽样都受到了随机变量的影响，但是它们受到影响的程度却是不一样的。第 2 个盒子的变化要小于第 1 个盒子的变化。这一点可以由第 2 个盒子中心值的窄幅聚集得到证明。我们看到第 2 个盒子的抽样变异水平比较低，所以我们对 $F-G_2$ 比 $F-G$ 的数值更具确定性的这一说法是恰当的。

我们从盒子的实验当中学到了哪些统计学知识

正如纽约扬基队前经理尤吉·贝拉 (Yogi Berra)^①所说：“只要去看就能观察到许多东西。”如果他是一位统计学家的话，他也许会说，通过抽样你就能够学到许多东西。即使一个抽样仅仅是一个更大的领域（总体）的一部分，它仍然可以使我们从中学到更多关于这个总体的东西。然而，一

① 前美国职棒大联盟的捕手、总教练，在 1972 年被选入棒球名人堂，被誉为“美国联盟最有价值球员”之一。——译者注

个抽样所告诉我们的内容是有它的局限性的，这一局限性是通过随机性来体现的。

统计学的一项基础任务就是证明由抽样变异性造成的不确定性。这就使得对基于抽样所做的表述的可靠性进行量化成为可能。如果没有这种量化，这些表述就如同受到限制的数值一样。统计学不能够消除不确定性，但是它可以告诉我们相关的信息，并且可以通过此举评估我们了解了多少以及我们了解的具体程度。因此，统计学分析对于过度自信的倾向来说是最有力的解决方法。

统计理论

盒子中的珠子这个实验涉及许多统计学推理当中的重要理念。此处我们加入一些其他理论来拓展这些理念，并且对技术分析法则进行评估。

统计学推理问题的 6 个因素

对诸如一种技术法则具有预测力、一种新的疫苗可以预防疾病或者其他知识的主张进行评估就是统计学推理问题的样本。一般说来，这些问题可以归纳为 6 点：①主体，②从主体当中随机挑选的包含一组观察资料的样本，③总体参数，④抽样统计，⑤推理，⑥对推理可靠性进行的陈述。当任意一种因素与“盒子中的珠子”这个实验相关联的时候，我们会对它们逐一进行讨论，然后，我们再对技术分析法则进行评估。

1. 主体

主体包含一个随机变量中所有可能的观察资料。主体很大（也许是无限大），但它是一个定义完整的观察资料主体。在典型的统计学推理问题中，我们需要学习一些有关主体的事实，但是对其进行全面的观察是不太实际，也是不可能的。在这个实验中，主体包括一组倒入盒子当中的珠子。颜色是随机变量。至于对技术分析法则的测试，主体包括所有根据该法则的信号在马上实现的预期中获得的每日收益率。¹³

术语马上实现的预期是什么意思呢？我们对处于运动当中的金融市场

是静止的假设是没有道理的,所以,对一种法则的盈利性是永远持续的预期也是不合理的。正因如此,主体对于技术分析法则来说,就不可能是指在无限的未来当中获得的收益率。

对于**马上实现的预期**,这里有一个更合理的解释。它是指一段有限的未来时间,在这段时间里,即使市场是不稳定的,但是对一个有价值的法则的盈利性可以持续的预期是合理的。任何致力于发现有预测力的形态的行为都必须对预测力的持续性做出假设。换句话说,除非有人愿意对预测力的持续性做出假设,否则所有的技术分析形式都是没有意义的。我们在这里所做的假设是用一种可以持续工作很长时间的法则来弥补研究者为了发现它而付出的努力。这种说法与格罗斯曼(Grossman)和斯蒂格利茨(Stiglitz)在“论信息效率市场的不可能性”(On the Impossibility of Informationally Efficient Markets)¹⁴一文中所说的相一致。这也与我们将在第7章讨论的主题“可以盈利的法则信号可以获得风险溢价的机会”¹⁵的说法是一致的。

因此,**马上实现的预期**是指在有限的未来时间里,所有可能的随机市场行为的实现。这就好比这里有一个无限的平行世界一样,除了市场行为的那部分,其他所有的世界都是完美的复制品。在每一个总体当中,对法则的盈利性做出解释的图形形态都是一样的,但是市场随机的那部分内容却是不一样的。这个观点如图4-12所示。

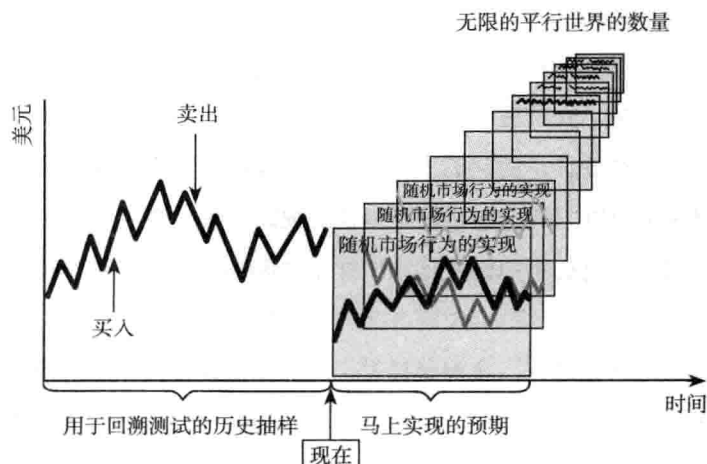


图 4-12 在无限的平行世界中，不同的随机实现

2. 样本

样本是可以用作观察资料的主体的子集。在“盒子里的珠子”这个实验中，我们可以观察 50 个独立的样本。对于技术分析法则来说，我们将有代表性地观察一个单独的样本，即通过回溯测试，来观察一种法则在市场历史的某一时段的绩效。这个样本包括一组由法则的信号所产生的连续的每日收益率。我们把这个收益率的数列减少到单个数字，即一个样本绩效的统计数字。

3. 总体参数

总体参数是指有关我们需要知道的这个主体的事实以及它的特征。总体参数是有代表性的数值，但是这里并不需要。

不幸的是，由于主体的全部并不能够被观察到，所以总体参数的值是未知的。统计学推理的本质就是试图加大有关总体参数的认知，尽管存在着对总体参数进行完全观察的不可能性。在“盒子里的珠子”这个实验中，总体参数的区间利益是盒子中灰色珠子的部分（即 $F-G$ 和 $F-G_2$ ）。而在技术分析法则的案例中，总体参数是指法则在“马上实现的预期”中的预期绩效。其绩效可以通过多种方法来确定。有一些最常用的计量方法，例如：平均收益率、夏普比率（Sharpe ratio）以及获利因子（profit factor）等。本书对绩效的计量采用的是基于消除市场趋势数据（零中心，zero-centered）中的年平均每日收益率。

在众多的统计学问题中，对总体参数是永远不会改变（即常量）的假设是比较安全的。在“盒子里的珠子”这个实验中，灰色珠子的比例保持着一个常量。在统计学的问题中，总体参数保持为一个常量是指稳定的意思。¹⁶ 现在试想以下的这种情况所包含的缺陷。假设在实验者不知情的情况下，一只看不见的手悄悄地把不同样本之间的灰色珠子进行增加或者减少，现在，总体参数，即 $F-G$ 是不稳定的，或者是非固定的。¹⁷

在本章的开始部分，我曾经说过，最好以假设任何接受检验的法则不具有预测力开始本章的内容。也就是说，我们假设该法则的预期收益率等于 0。或者，用统计学的术语来说，就是总体参数假设等于 0。

4. 抽样统计

抽样统计是指一个样本的可测量属性。¹⁸ 因为该样本已经被观测过，所以它的数值是已知的。本书中，术语**抽样统计**是指用数字证明的事实。例如，比率、百分比、标准差、平均收益率、裁剪平均值¹⁹以及夏普比率等。在统计学的推理问题中，**抽样统计**特指和总体参数一样的可测量属性。在“盒子里的珠子”这个实验中，抽样统计，即 $f-g$ ，是指在独立的样本中灰色珠子的数量。而总体参数，即 $F-G$ ，是指整个盒子当中灰色珠子的数量。

从本质上来说，抽样统计因为揭示了总体参数，所以非常重要。除此之外，它没有其他特别重要的事实。有些市场历史学家看似并不留意这些重要的事实。

如果回溯测试的结果是正收益率，那么就会产生以下的问题：回溯测试产生了一个大于0的随机偏差正值绩效，这到底是因为抽样变异性，还是因为该法则具有预测力（该法则的预期收益率大于0）呢？回答这个问题需要运用统计学推理的工具（见图4-13）。

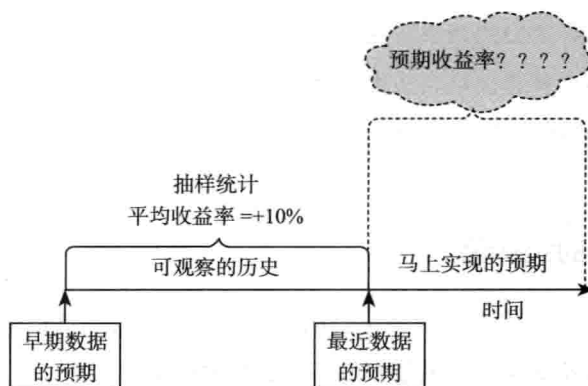


图 4-13 参数的数值真的大于 0 吗

5. 推理

统计学推理是从被观察的抽样统计值中形成的感应式飞跃，它以确认为人们所知，但是仅仅对于特定的数据样本来说才是真实的。对于总体参数的数值来说，这个值是不确定的，但是对于一个范围广的、无限的以

及没有被观察到的情况的数量而言，这一数值又看似是真实的。

当一种法则过去的正值绩效可以合理地归结为抽样变异性的时候，则其对马上实现的预期所做出的预期收益率是 0 或者小于 0 的推理就是合理的。然而，如果绩效的正值太高，而不能合理地归结于抽样变异性（运气）的时候，合理的推理就应该是这样的：该法则具有真正的预测力，并且认为马上实现的预期收益率为正值。

6. 对于推理可靠性的陈述

由于一个抽样并不能代表总体，基于抽样统计的推理就会容易产生不确定性。换句话说，这种推理有可能是错误的。我们已经看到抽样统计是如何围绕着总量参数的真实数值进行随机的变化了。科学的统计学已经超越了承认其推理是模糊不清的境界。它把其可靠性进行了量化。这就使得统计学的结论要比那些通过不能够提供信息的非正规的方法所得出的结论更有价值。

在两种情况下推理会导致错误。第一，当法则不具有预测力的时候得出其具有预测力的结论。这种情况属于好运气降临到没有价值的法则身上的范畴。把这种类型的错误转变成为对市场风险的承受能力，其后果是无法弥补的。第二种错误是当法则具有真实的预测力时，却得出其不具有预测力的结论。这种错误会导致错过交易机会。

✓ 描述性的统计学

统计学领域可以分为两个主要的部分：描述性的和可以进行推论的。对于技术分析法则来说，最重要的事情就是进行统计推断，这一点我们已经在前面的部分中讨论过了。然后，在进行推理之前，我们必须用简明的、信息含量丰富的方法对抽样数据进行描述。而描述性的统计学正好适用于这个目的。

描述性的统计学的目的就是进行数据简化。也就是说，把一组大的被观察数值简化成为规模较小的数据，使之成为更加明白易懂的一组数字和图形。描述性的统计学讲述的是整体的故事，而不是单独个体的故事。有

3种描述性的工具对于摆在我们面前的任务是非常重要的: ①频率分布, ②对集中趋势的衡量, ③对变异的衡量。

频率分布

我们在之前“盒子里的珠子”这个实验中已经讨论过频率分布了。它们把用于随机变量 $f-g$ 的一组包含 50 个观察资料的数据简化成为一组信息含量丰富的概要。

首先, 由于只有 50 个观察资料, 我们可以把它们组成一个看上去由一组数字的数据整体。其次, 如果观察资料的数量是 500 个或者 5 000 个, 则这组数字就未必像频率分布那样含有丰富的信息了。

频率分布通常是分析一组数据的第 1 步。它快速、直观地表现出一个观察资料样本的两个主要特征: 集中趋势(例如, 平均值)和中心值的变化或者偏离程度。例如, 图 4-11 所示的两个盒子当中珠子的中心值和偏离程度都不尽相同。

一个抽样的效果是有价值的, 但是量变可以产生质变。为了我们的研究目的, 我们需要把频率分布的两个特征进行量化: 集中趋势和它的变异性, 或者集中趋势的分散情况。

对集中趋势的衡量

对集中趋势的衡量有很多种方法。其中 3 种最为常见的方法是平均值、中值和众数。平均值, 也称为算术平均值, 应用于许多的技术分析当中, 也是本书中所采用的概要统计。它是用观察资料的数量除以被观察的数值的和。

从抽样的中值当中分辨出总体的中值是非常重要的。抽样的中值是一个已知的统计数字, 它是从连续的抽样当中观察到的数值与随机变量计算得出的。而总体的中值则与抽样的中值相反, 它是在固定的问题当中存在的未知的、非随机的数值。抽样的中值的公式有两种(见图 4-14)。

变量 X 的抽样中值

$$\bar{X} = \frac{x_1 + x_2 + x_3 + \cdots + x_n}{n}$$
$$\bar{X} = \frac{\sum_{i=1}^n x_i}{n}$$

其中, x_i 代表变量 X 中的独立的观察资料

图 4-14 变量 X 的抽样中值

统计学中的变异（分散）测量

变异测量是指在抽样当中，观察资料与它的集中趋势所产生的离差的程度。换句话说，变异测量把频率分布的宽度进行了量化。

最为广泛应用的衡量离差的工具包括：方差（variance）、平均绝对偏差（mean absolute deviation）以及标准差（standard deviation）。它们是传统的统计学分析方法（经典统计法）中最重要的 3 种方法。标准差是数据中值中每一个观察资料的平均方差的平方根。一个抽样的标准差可以用图 4-15 中的公式来表示。

观察资料抽样的标准差

$$s = \sqrt{\frac{\sum (X_i - \bar{X})^2}{n}}$$

其中， X_i 表示变量 X 中独立的观察资料； $\bar{X}$ 表示变量 X 中抽样的中值； n 表示观察资料的数量

图 4-15 观察资料抽样的标准差

有关离差的更为直观的概念可以用图画的方式来表示。图 4-16 显示了一组由不同的变异程度和集中趋势所组成的理想的频率分布。它们在某种意义上被认为是理想的，是因为图中具有实际频率分布特征的阶梯状已经被消除了。图 4-16 的关键点在于集中趋势和变异（离差）具有频率分布的独立特征。在图 4-16 的第 1 列中，4 个分布具有相同的变异程度，只是中

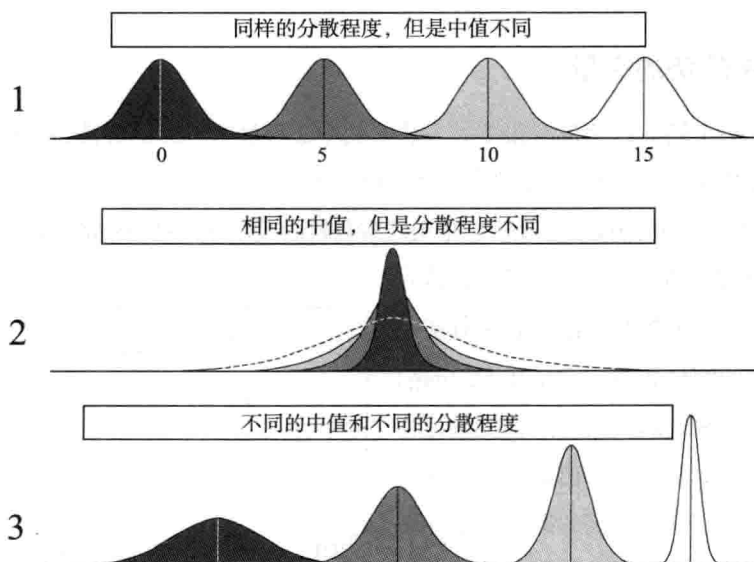


图 4-16 集中趋势和分散是分布的不同属性

心值不同而已。第2列中,分布具有相同的中心值,但是变异的程度却不相同。第3列中所有的分布都有不同的中心值和变异程度。

概率

概率这个概念在统计学中是非常重要的,因为它将不确定性进行了量化。从统计学推理当中得到的结论具有不确定性;也就是说,有可能会出现错误。因此,对总体参数的值所做的推理在一定程度上是错误的。我们把这种可能性用术语“概率”来表示。

我们按照惯例用表示概率的非正式概念来形成预测并做出选择。我遇到的这个人现在是如此的漂亮,结婚之后还会如此吗?如果我在一个特定的地点不停地挖,那么我挖到金子的可能性有多大?如果我从事律师行业,实现我的志向的概率有多大?我到底该不该买入那只股票,如果我买了,那么我能赚取多少利润呢?

我们把对概率一词的非正式使用与图4-17中所示的可以互换的术语的子集联系在一起。

为了达到这个目的,我们需要获得一些更具确定性的东西。我们要对基于相对频率的概念所要面对的工作做出对概率

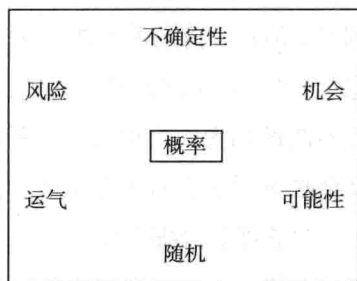


图 4-17 概率的一般概念

这一概念最好的解释。某一事件的相对频率是指用事件可能发生的概率的次数除以事件实际发生的次数计算得出的。相对频率的值被设定在 0 ~ 1。也许现在天天都下雨,但是在 4 月下雨的相对频率是 0.366 (11/30)。如果降雨从来都没有发生过,那么相对频率就应该是 0。如果每天都下雨,则相对频率就是 1.0。因此,相对频率的结果就等于

$$\text{事件发生的数量} / \text{事件可能发生的数量}$$

在“盒子里的珠子”这个实验中, f-g 的值被测试了 50 次。50 这个值对于任何出现的 f-g 值来说是最大值了。实际上, f-g 的值等于 0.65,也就是说在 50 个 f-g 值当中, 0.65 这个值一共出现了 5 次。因此,它的相对频

率就是 0.10 (即 $5/50$)。

一个事件在很长的时间里——非常长的时间里出现的相对频率可以用概率这个词来表示。也就是说,一个事件的概率就是它出现的相对频率,即对事件的发生给出一个无限的数量。概率被描述成为一个范围位于 $0 \sim 1$ 的数字。概率为 0 意味着该事件从未发生,而概率为 1.0 则意味着该事件经常发生。

概率是一个理论上的概念,因为我们永远不可能去观察任何一个无限的数字。然而,出于实践的目的,当观察资料的数量变得非常大的时候,相对频率就会向理论上的概率靠近。

大数法则

当观察资料的数量变大时,相对频率的趋势向着理论上的概率这一点集中,这就是大数法则。²⁰ 至于大数法则的应用,我们可以用投掷硬币的例子来演示。掷硬币可能出现的结果就是出现正面或者背面,我们假设头像(正面)的出现概率是 0.50。大数法则要告诉我们的是随着掷硬币的次数不断的加大,出现头像的相对频率会朝着理论上的值 0.5 集中。然而,即使掷硬币的次数不断加大,但是偏离 0.50 这个概率值的情况仍然会发生,尽管随着投掷次数的上升,可能出现的背离规模会下降。然而对于小样本来说,大数法则提醒我们出现头像(正面)的相对频率可能会大幅偏离 0.50。当掷硬币的次数只有 3 次时,1.0 这个值会很容易出现,即 3 次投掷都出现头像(正面)。

图 4-18 显示了在掷硬币实验中,投掷的次数由 1 次到 1 000 次出现头像的部分。当掷硬币的次数(观察资料)小于 10 次时,珠子在 0.50 这个点经历了两次较大的偏离。在第 3 次投掷的时候,珠子的比率为 0.66 ($2/3$)。在第 8 次投掷的时候,这个比率为 0.375 ($3/8$)。然而,随着投掷次数的增加,随机变量、出现头像的部分在 0.50 这个预期的数值上所经历的偏离越来越小了。这清楚地表达了大数法则所预测的内容。

现在,我们来试想一个天真的掷硬币者。他在掷出的 5 次硬币中,出现了 5 次头像(正面),然后他大声地宣布:“我找到了硬币中的圣杯,它

总是出现头像这一面。”而大数法则会告诉我们这个可怜的人的乐观态度是毫无根据的。这就好比技术分析的研究者发现了5个历史性的信号，而全部信号都被证明是正确的一样。乐观主义情绪是毫无根据的。在第6章中我们会说明如果这种法则被选择，是因为它在大量的进行回溯测试的法则中所表现出的优异绩效，也就是说，它是通过数据挖掘被发现的，所以对于该法则在未来所表现出的绩效的乐观情绪也就大大地降低了。

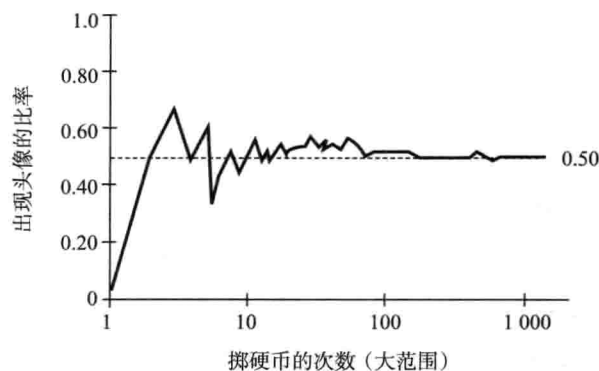


图 4-18 大数法则

理论概率 VS. 经验概率

概率分为两种类型：理论上的和经验上的。“理论概率是由基于逻辑基础的纯粹的自信程度所决定的。它们是可以从之前的经验中推导出来的，而其中的大部分是以对称性论证法为基础的”。²¹ 换句话说，它们可以不需要对经验的证明就可以推导出来。在掷硬币的实验中出现头像的概率就属于理论上的概率。我们并不需要掷1万次硬币来预测在漫长的掷硬币过程中出现头像的概率是0.50。简单的事实是硬币是公平的，因为它有正反两面，我们不可能在硬币落在其边缘的时候就有充足的理由来猜测正面与反面是相等的。诸如在发扑克牌的时候出现同花顺、在掷色子时出现6点以及赢得彩票的机会的概率都属于理论概率，因为它们都可以通过对情况的逻辑分析来决定。

经验概率是以对频率的观察为基础的。技术分析就涉及这种类型的概率。7月胡德山下雪的概率，或者说随着股票市场的上涨而来的是利率下

跌 2% 的概率就是经验概率的实例。我们只有对过去特定条件（7 月的胡德山）下的实例进行大量的观察才能对此做出决定，并且确定该事件（新的降雪）的相对频率。

当确定经验概率的时候，通过相同的一组条件确定每一个实例的特征是非常重要的。在技术分析领域里，这种做法是绝对不可能成为现实的。过去每一个利率下跌 2% 的实例都与那个特定的条件相类似，但还有许多条件不同。例如，在利率水平下跌 2% 以前，每一个观察资料的情况都是不同的。在一个实例当中，利率在下跌之前处于 5% 的水平，而在另一个实例当中，利率在下跌之前则处于 10% 的水平。当然，利率的水平也要加入到定义每个实例的各组条件当中去，但是这样做会降低可以进行比较的实例的数量。另外，还存在不属于特定的条件设置范围的其他情况。在非实验科学和观察科学当中，不能够控制所有的潜在相关变量是不争的事实。反之，实验科学家可以享受到除正在被观察的变量之外，控制其他所有常量的有利条件。想象一下这种情况出现在技术分析当中会是什么样的情形！对此，我想我只有在来世的时候才能解释了。

随机变量的分布概率

概率分布揭示了我们期望不同的随机变量可能出现（它们的相对频率）的数值发生的频率。随机变量的概率分布是从无数的观察资料中所建立的相对频率分布。

概率分布这个概念可以通过考虑相对频率分布的排列顺序的方式来理解，每一个频率分布都是从不断增加的观察资料数量中建立的，并且逐步发展成为更加狭窄的箱体或者区间。相对频率分布逐渐地成为概率分布。随着观察资料数量的增加，以及区间宽度的不断缩小，更多不相关联的区间就会产生。²² 由于观察资料的数量接近无穷大，导致了区间的数量也接近无穷大，并且使得它们的区间宽度减至零。因此，相对频率分布就会转变成为我们所说的概率分布，从更加专业的角度来说，它指的是概率密度函数。²³

图 4-19 ~ 图 4-23 显示了在区间数量增加, 并且区间密度不断缩小时, 相对频率分布是如何演变成为概率密度函数的。每一个柱子的高度代表着落入该柱状图区间的事件的相对频率。这个数字假设随机变量代表了价格变动的程度。

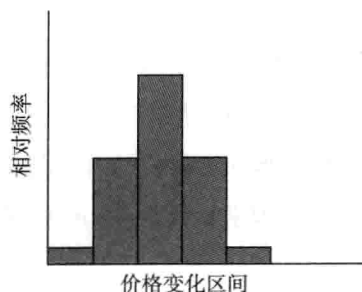


图 4-19 基于 5 个区间的
相对频率分布

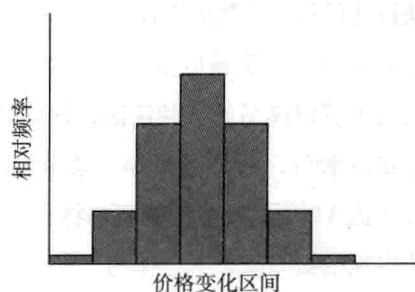


图 4-20 基于 7 个区间的
相对频率分布

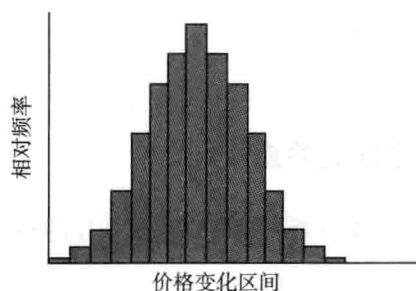


图 4-21 基于 15 个区间的
相对频率分布

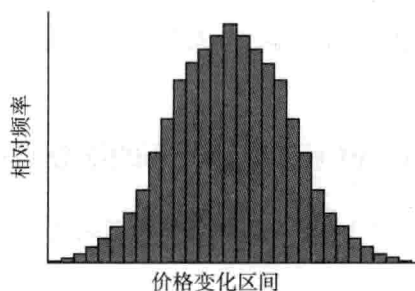


图 4-22 基于 29 个区间的
相对频率分布

这组连续的图形揭示了概率密度函数实际上就是包含无限观察资料, 并且其区间宽度是无限狭窄的相对频率分布。但是事情有些奇怪之处。如果概率分布区间的宽度为零, 那么每个区间里的观察资料也都为零。这个说法看上去是毫无意义的! 然而, 对于数学概念来说, 其与常识之间存在的矛盾是非常正常的。在几何学当中,

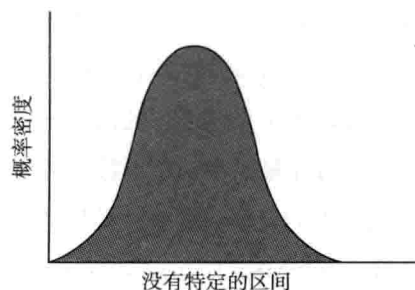


图 4-23 概率密度函数

一个点对应一个定位,但是它并没有占据任何空间(例如,长度、宽度以及广度等于零),例如一条线只有长度而宽度为零等。

包括零观察资料的区间具有一种奇怪的含义,即任何连续的随机变量出现单一值的概率等于零。出于这一原因,只有在提及这种概率时,即该随机变量将一个数值假设在一个特定的数值范围内,这种说法才是合理的。换句话说,就是随机变量将一个数值假设在特定的最小值与最大值之间。我们也可以换另外一种说法,即将随机变量的数值假设等于或者大于某个特定的数值,或者小于等于某个特定的数值。例如,我们可以说出现 3.0 或者比 3.0 更大的数值的概率将会发生。然而,如果我们提出 3.0 这个精确值出现的概率就是不合理的。²⁴

这种有违直觉的思想非常符合对法则的测试。当测试一个法则过去收益率的统计学意义时,我们就会涉及收益率在 +10% 的水平以及在更高收益率水平时出现的概率,前提是该法则是在不具有预测力的条件下依靠运气的成分获得的。概率密度函数就可以证明这些数据。

概率和概率分布的部分面积之间的关系

(1) 某一事件的概率就是在无限的观察资料中该事件将要发生的相对频率。

(2) 概率密度函数是从无限的观察资料和宽度为零的区间中建立的对频率分布。

(3) 随机变量的相对频率呈现出一个在给定区间内的数值,这个数值与落入该区间顶部的频率分布的面积是相等的(见图 4-24)。

我们接下来要进行的最后一步与上面所列的第 3 点类似,只是术语相对频率分布被概率密度函数所替代。连续的随机变量假设该数值出现在特定范围内的概率与围绕在范围周边(之上)的概率密度函数部分相等。这一概念如图 4-25 所示。图 4-25 显示了一个连续的随机变量 X , 其假设在 A 与 B 之间的数值等于 0.70。

在许多的实例当中所涉及的区间是指分布的极端值或者尾数。例如,

图 4-26 显示了随机变量 X 的概率，它假设 B 的数值相对更大，即概率是 0.15。

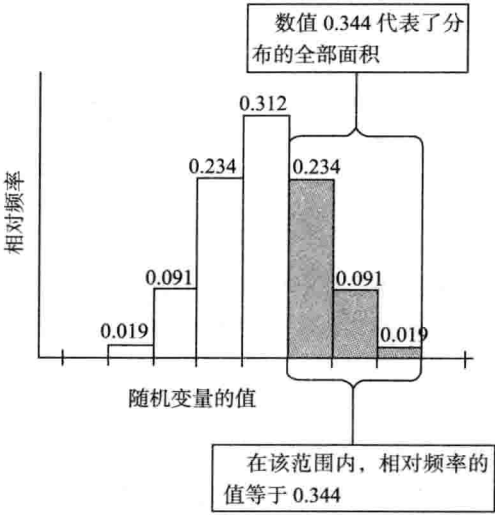


图 4-24 相对频率与分布的部分面积相等

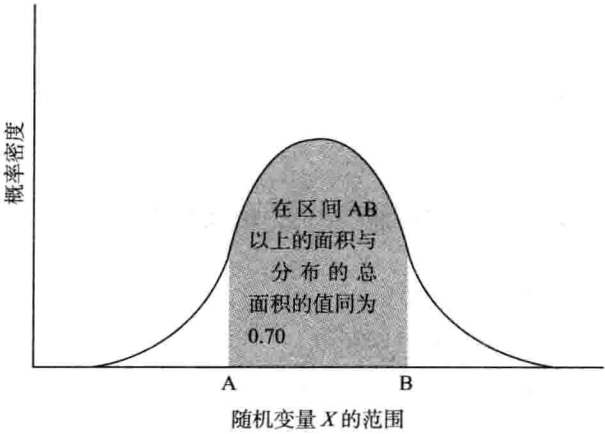


图 4-25 概率随机变量 X 假设 A 与 B 之间的值等于在区间 A 与 B 之上的分布的全部面积

随机变量的概率分布是个非常有用的概念。即使对随机变量 X 的未来进行独立的观察是不可预测的，但是大量的观察资料会组成一个具有高度预测性的形态。这一形态就是随机变量的概率密度函数。

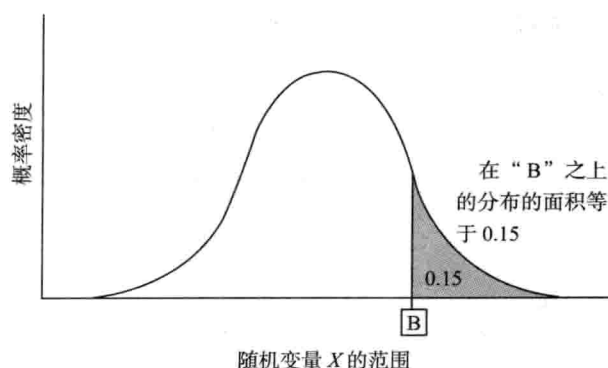


图 4-26 概率随机变量 X 假设一个等于或者大于 B 的值
与等于或者大于 B 的分布的全部面积是相等的

现在，就是这些想法与我们见面的时刻了。第 5 章将涉及当一个假设由随机变量的概率分布所代表的时候，我们可能会运用一个被观察到的随机变量的数值来测试该假设的真实性。这样一来，我们就可以得出一个法则在回溯测试中产生的盈利能力到底是由于运气的成分还是由于其本身所具备的真正的预测力所导致的结论了。

用于这一目的的概率分布是一种特殊的类型。我们将其称为抽样分布——也许它在所有的统计学和技术分析师当中是最为重要的了。

✓ 抽样分布：统计学推理当中最重要的概念

回顾一下已经学过的知识。

(1) 随机变量是概率实验的结果。

(2) 概率实验²⁵就是指从总体中提取一个给定的观察资料数值的样本，然后计算诸如抽样的平均值这样的统计数据。

(3) 抽样统计²⁶是指在不同的抽样之间进行的没有预测性的波动的随机变量。抽样统计之所以会出现随机的波动，是因为偶然性决定了在一个特定的抽样当中特殊的观察资料的结果，正是这组特定的观察资料决定了该抽样的统计学数值。在“盒子里的珠子”这个实验中， $f-g$ 的值在不同的包含 20 个珠子的实验中呈随机变动的状态，这是因为一组特定的珠子在这

个给定的抽样当中的结果是由偶然性决定的。

(4) 相对频率分布描述了在大量的观察资料中, 随机变量的数值可能出现的频率。

我们现在来讨论抽样分布这一在统计学推理当中最为重要的概念。

抽样分布的定义

抽样分布 (sampling distribution) 就是随机变量的概率分布, 而这个随机变量恰好就是样本统计量 (sample statistic)。换言之, 抽样分布显示了我们可以假设到的各种样本统计量的数值, 以及它们与概率的联系。²⁷ 例如, “样本平均值的抽样分布是指在总体²⁸ 中给定范围内的所有可能的随机样本的所有可能的概率分布的平均值”。因此, 样本统计量就是平均值。

在“盒子中的珠子”这个实验中的样本统计量就是 $f \cdot g$ 。它的随机变动横跨 50 个样本, 每一个样本包含 20 个珠子, 我们将其用相对频率分布的形式在图 4-8 中表示。我们用所有可能的 20 个珠子的样本代替 50 个数值——这个数量会非常庞大, 这种理论上的分布就是 $f \cdot g$ 的抽样分布。

在法则的回溯测试当中, 样本统计量就是在回溯测试中观察到的衡量绩效。本书中的样本统计量指的是法则的平均收益率。由于我们只有一个单一的市场历史样本, 所以回溯测试对绩效的统计量只产生了单一的观察值。现在想象一下如果我们可以市场历史的无限多的独立样本中进行回溯测试时的情景吧。这就会为绩效统计量提供无穷大的数值。如果这组数据随后转变成为相对频率分布的话, 那么它就是精确的统计学抽样分布。很显然, 这一点是不可能实现的, 因为我们并不拥有无穷大的独立的历史数据样本。然而, 统计学家发明了很多种获得接近精确的抽样分布的方法, 尽管实际上我们只有一个历史数据的样本和一个抽样统计量的数值。我们在本章接下来的部分里将要介绍两种主要的方法。

抽样分布对不确定性进行量化

统计学中的抽样分布是统计学推理的基础, 这是因为它把由抽样的随机性 (抽样变异性) 所导致的不确定性进行了量化。

正如我们在前面所表述的那样，抽样分布表示了如果它在从总群中抽取的，且同样规模中的无限数量的随机样本中进行测试时，统计数据的相对频率。“盒子中的珠子”这个实验显示了统计数据值， $f-g$ 的值在不同的样本之间的波动是随机的。图 4-8 显示了这些数值倾向于落入一个表现良好的形态，而不是一个随机混乱的图形。实际上，表现良好的随机变量的形态更加有助于统计学的推理。

具有讽刺意味的是，样本统计量显示，实际上是一个随机变量产生了这样一个有规律的形态。我们要把这一切都归功于统计学家们。图 4-8 显示了一个大约 0.55 的中心值（抽样分布的平均值）。而且它还显示了有关这个中心值的定义完好的离差形态。这一形态使得我们可以得到这样一个结论，即在有一定把握的前提下，总体参数 $F-G$ 的值保持在 0.40 ~ 0.65 的范围内。从中我们还可以得到一个结论，就是尽管在把握性降低的情况下， $F-G$ 的值会更加精确地落入 0.50 ~ 0.60 这个范围。正是因为 $f-g$ 的抽样分布中离差的程度，才使得它有可能对总体参数 $F-G$ 的数值做出表述。

抽样分布的离差对我们有关总体参数 $F-G$ 知识的不确定性进行了量化。分布的集中趋势传递着最有可能是 $F-G$ 值的信息，即 0.55。知道了这些固然非常好，但是还远远不够。更重要的是我们要知道 0.55 这个数值的真实性。换言之，即用数值 0.55 这个抽样分布的集中趋势来描述 $F-G$ 值的真实性的准确度有多大。

这一准确度（确定程度）是通过抽样分布的离差来表达的。围绕着中心数值 0.55 的分布的离差越大，那么由数值 0.55 确认的在整个盒子中灰色珠子的比例 $F-G$ 数值的真实性的精确程度就越小。

为了显示这一点，我们来分析下面两个抽样分布。这两个抽样分布的中心值都是 0.55，但是它们的离差值却各不相同。图 4-27 中显示的第 1 个抽样分布的离差较小，因此它给我们以强烈的印象来说明 $F-G$ 的值就在 0.55 的周围。而在图 4-28 中所显示的离差分布的幅度很大，以至于它不能够提供更多关于 $F-G$ 的数值的信息。这就是说真正的 $F-G$ 值也许和抽样分布的中心值存在很大的不同。

总的来说，确定性与抽样分布的离差直接相关。而关于总体参数的结

论的确定性则依靠于统计数据的抽样分布的宽度——宽度越大意味着不确定性越强。

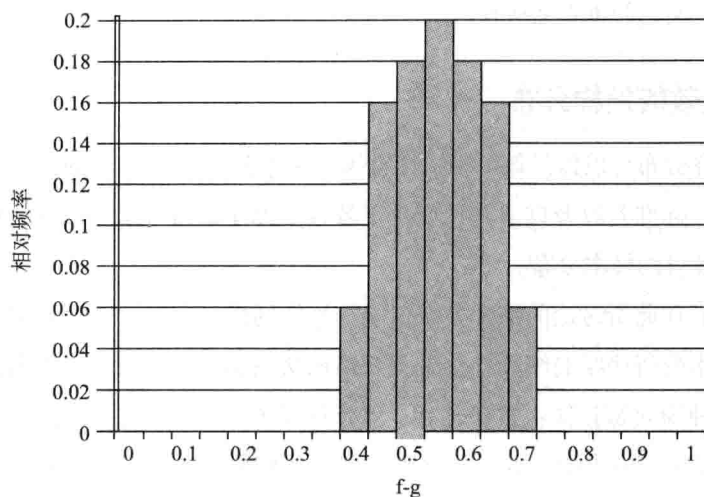


图 4-27 相对频率分布: $f-g$ (样本规模=20, 样本数量=50)

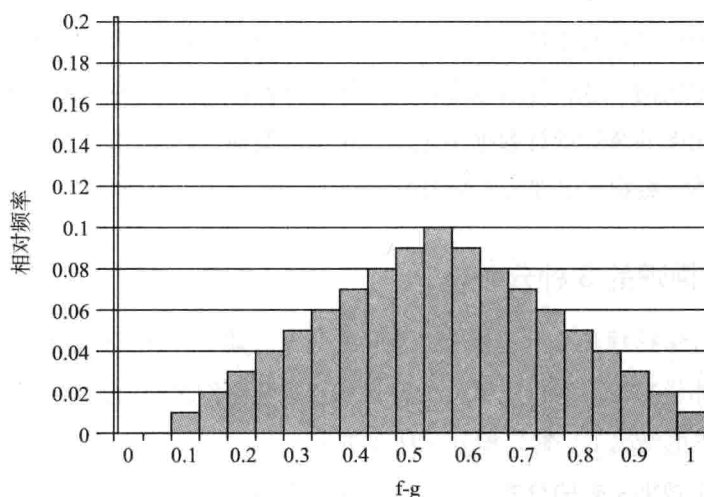


图 4-28 相对频率分布: $f-g$ (样本规模=20, 样本数量=50)

假设一个法则的预期收益率等于零。再假设该法则在回溯测试中的收益率大于零。那么我们可以说正的收益率足以证明收益率为零的假设就是错误的吗? 对这一问题的回答要依靠相对于抽样分布的宽度而言, 回溯测试的收益率超过零的距离来确定。如果相对偏差较大, 那么这个假设就被

否定了。相对偏差的量化与是否放弃这种假设的决定背后的推理，我们将会在第5章中进行讨论。到目前为止最重要的一点就是抽样分布的宽度是回答这一问题的重点之所在。

交易绩效的抽样分布

抽样分布可以以任何一种统计量的形式存在：平均值（mean，平均值）、中位数、标准差以及许多其他应用于统计推断中的统计量。²⁹每一种统计量都有各自的样本分布。

我们在此所讨论的抽样分布仅限于平均的抽样分布，因为本书中用于评估技术分析法则的绩效统计量就是指该法则的平均收益率。然而，还有其他多种绩效衡量法：夏普比率、³⁰ 获利因子、³¹ 平均收益率除以溃疡指数、³² 等等。应当指出的是，这些可以替代的绩效统计量的抽样分布与平均值的抽样统计量是不同的。

我们在此还需要指出的是，本书中所应用的产生平均值的抽样分布的方法，在用延长右侧尾部的方法来产生绩效统计量的抽样分布中，它的值是受到限制的。这在包含夏普比率、从平均收益率到溃疡指数以及获利因子这些比率的绩效统计数据中是会发生的。然而，这个问题可以通过用比率的对数缩短右侧尾部的方法来缓解。

统计学推理的3种分布

统计学推理实际上包含3种不同的分布，其中一种就是抽样分布。它们之间非常容易混淆。这部分内容就是要区分它们之间的差异。在对法则进行测试的情况下，这3种分布具体如下。

（1）**数据总体的分布。**一个规模无限大的分布包括所有可能得到的该法则每日的收益率，我们把对它的假设延伸到马上实现的预期这一范围。

（2）**数据抽样的分布。**包括该法则从过去到现在的数量（N）有限的每日收益率。

（3）**抽样分布。**规模无限的样本统计量分布——这里指的是法则的平均收益率，如果将法则在从总体中提取出来的一个拥有无限随机样本的规

模 N 中进行测试的话, 那么它就代表着该法则的平均收益率。

有两点我们需要强调一下。首先, 观察资料是由包含总体的数据分布和抽样的数据分布 (上面所列清单中的 (1) 和 (2)) 所组成, 它们都是由法则每一天的收益率组成的。反过来, 包含抽样分布的观察资料就是样本的统计量, 其中, 每一个观察资料代表着通过对数天的样本进行计算得出的该法则的平均收益率。其次, 总体分布和抽样分布在这种情况下都是理论性的, 它们所指的是一个拥有无限数量的观察资料。反之, 样本的抽样分布包括在历史回溯测试中数量有限的观察资料。

3 种分布的差异可以通过假设一个与“盒子里的珠子”相类似的实验便可以直观地看到。然而, 在这个例子中, 你可以想象在规模上是无穷大的每日收益率的总体分布。设想我们采用 50 个从总体当中提取的独立样本, 每一个样本包括大量的每日法则的收益率。接下来要决定每一个样本的平均收益率, 然后绘制 50 个平均值的抽样分布 (见图 4-29)。³³

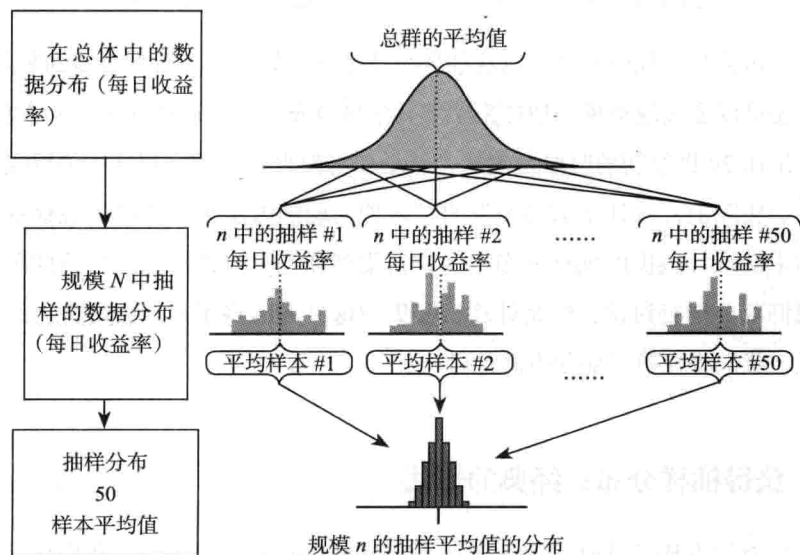


图 4-29 统计学推理的 3 种分布

真实的世界：单一样本的问题

我们之前假设了如果从总群当中提取 50 个独立抽样进行观察时将会出

现的情况，我们把这种讨论称为理论上的。在大多数现实世界中的统计学问题中，通常会遇到只有一个观察资料的样本，且它的平均值也只有一个。问题是在只能获得一个抽样平均值的情况下，我们不可能掌握抽样统计量的变异性。一个样本与一个平均值的问题如图 4-30 所示。

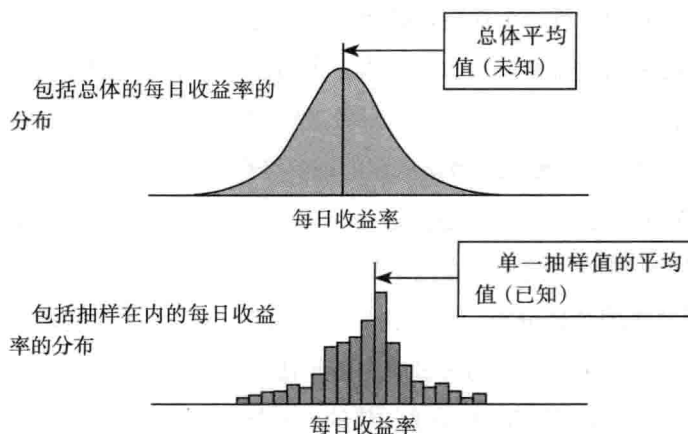


图 4-30 现实的世界：单一样本与检验统计量的单一值

幸运的是，我们不仅可以从抽样统计量的变异性当中学到很多东西，而且还可以通过观察单一抽样来了解其抽样分布。至于做这项工作的方法，最早是在 20 世纪初的时候发现的，就是如何使得统计学推理这一领域成为可能。实际上，统计学家已经发明了两种明确的方法来估算单一观察资料样本中的统计数据的抽样分布：经典密集型和计算机密集型。尽管两种方法我们都要进行讨论，但是计算机密集型这种方法将用于对法则进行绩效测试，我们将在第二部分中进行讨论。

获得抽样分布：经典的方法

经典的方法是基础的统计学课程中通常讲授的课程之一。我们认为这主要归功于罗纳德·费希尔（Ronald Fisher，1890—1962）爵士和杰兹·内曼（Jerzy Neyman，1894—1981）这两位数理统计学之父。这种方法在被观察的单一样本的基础上，利用概率理论和积分学来获得抽样分布。它提供了对抽样分布的离差、平均值及其基本形态（普通）的估计值。换言之，

它提供了以单一观察资料样本为基础做出的推理的可靠性进行量化的所有需要的东西。

样本平均值的抽样分布

任何统计数据都有各自的抽样分布。我们对此的讨论将集中于平均值(平均)的抽样分布,因为它是我们将在第二部分中对技术分析法则进行评估时所要用到的统计数据。

经典的统计学理论会告诉我们很多与平均值的抽样分布相关的事情。

(1) 大样本的平均值可以对从总体分布中获得的抽样的平均值进行良好的评估。样本的规模越大,抽样的平均值就会越靠近总体的平均值。

(2) 抽样分布的离差取决于样本规模的大小。对于给定的总体来说,样本的规模越大,抽样分布就越狭窄。

(3) 抽样分布的离差也取决于总体数据内变量的总量。在总体中的变量越多,抽样分布就越宽广。

(4) 在将其应用于对技术分析法则的评估的情况下,随着抽样规模的增加对普通形态形成的确认,平均值的抽样分布形态将会向所谓的普通或者钟形的形态发展。而其他像比率这样的抽样统计数据,则会呈强烈的不规则形态。

对以上各种概念的解释如下。

1. 大样本的平均值可以对总体的平均值进行良好的评估

这是大数法则的表现形式,即如我们在之前掷硬币的实验中所见到的那样。我们现在来回顾一下,大数法则告诉我们包含样本的观察资料的数量越大,则样本的平均值就越会靠近总体的平均值。有一些符合这种观点条件,³⁴但是在此处我们不会去关注这些。

掷硬币实验的曲线图显示了大数法则(见图 4-18)的作用。随着掷硬币次数的增加,正面(头像)出现的概率就会向其理论上正确的值 0.50 集中。在实验的早期阶段,当样本的规模很小时,这时的数值与 0.50 出现了较大的背离。这些背离的数值显示了在小样本中运气的成分起到了很大的作用。在投掷 4 次的实验中,尽管出现 0.75 或者 0.25 的概率并不是最大

的,但是它们出现得非常频繁。然而,当投掷硬币的次数达到了60次,出现0.75或者0.25的概率则小于千分之一。我们从中获得的最重要的经验是:样本规模的扩大消除了运气的成分。

2. 抽样分布的离差取决于样本的规模:样本越大,离差越小

试想一个很大的观察资料的总体,其变量的标准差为100。通常情况下,我们是不知道总体的标准差的。但是在这个案例当中,我们假设这个值是已知的。

在图4-31中,我显示了在底部的总体分布。上面的两个抽样分布分别是样本规模为10和100的变量的平均值。请注意样本分布的宽度每次被切掉1/3的面积,而样本的规模则以10倍的速度增加。经典的统计学中有这样一种说法,即平均值抽样分布的标准差与样本规模的平方根成反比。³⁵因此,如果样本的规模以10倍的速度增加,那么抽样分布的宽度就减少了3.16倍(即约为10的平方根)。总体就可以被视为只有一个(单一平均值——总体平均数)样本规模的抽样分布。样本规模为100的抽样分布就是总体宽度的1/10。我们从中得到了非常重要的信息,即样本规模越大,样本统计量的值所存在的不确定性就越小。顺便说一句,平均值抽样分布的标准差在统计学当中有一个特殊的名字——平均数标准误差(请参见后面的解释)。

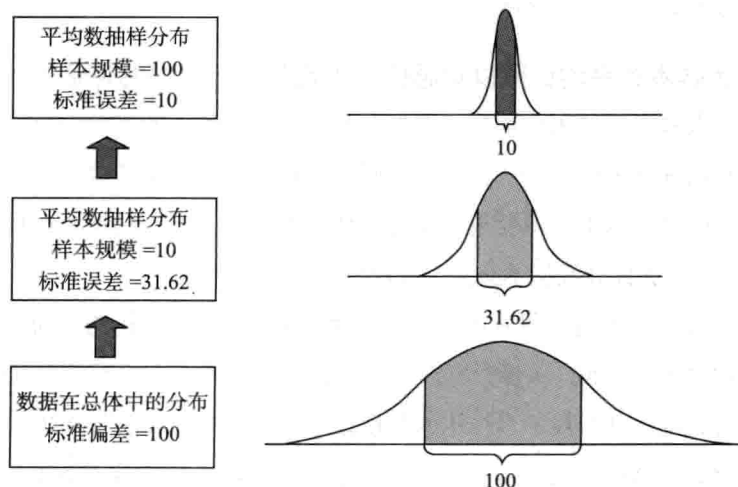


图 4-31 随着样本规模的增加, 抽样分布越来越狭窄

根据某些条件,³⁶大数法则告诉我们样本的规模越大,我们有关总体平均值的知识就越准确。但是需要注意的是,³⁷我们希望用尽可能大的样本来最小化抽样分布的宽度。

3. 平均数抽样分布的偏差也取决于包含总体在内的变量的总数

还有另外一个影响抽样分布的因素,即从总体当中提取的样本中包含的变量的总数。总体数据中包括的变量越大,抽样分布的变动(偏差)也就越大。

你可以通过总体中不包含变量这个想法来看待这个问题。例如,如果总体中的所有人体重 150 磅,那么每一个样本的平均重量就会是 150 磅。因此,如果总体的数量中不包含变量,那么就会导致样本平均值的变动为零,所以抽样分布的结果就不可能有偏差。反之,如果在包含总体在内的个体中出现较大的变动,那么样本的平均值的变动就会很大,并导致规模更大的抽样分布的出现。

4. 趋于正常抽样分布的形态

中央极限定理(Central Limit Theorem)是经典统计法的基础原理,即随着样本规模的不断扩大,在某些条件下,³⁸不论总体分布的形态如何,它都会向一个特定的形态发展。换言之,不论在总体中的数据的分布形态有多么的特别,抽样分布的形态都会形成一个特殊的形态,统计学家将其称为正态分布(normal distribution),或者称作高斯分布(Gaussian distribution)。

正态分布是最常见的概率分布形态。通常情况下指钟形曲线,因为它看上去就像是自由钟(Liberty Bell)的阴影部分,正态分布具有连续变量的特点,它可以用来解释很多现实世界中的现象。

对正态分布的理解,有几点需要我们注意。

- ▶ 正态分布由它的平均值和标准差来描述。如果这两个因素是已知的,那么你就会了解分布的具体情况。
- ▶ 大约 68% 的观察资料会落入单一的平均数的标准差中,而有 95% 的观察资料会落入两个平均数的标准差中。

- 在两个标准差之外的分布尾部会非常细小，因此，超过3个标准差的数值就会很少出现，而超过4个标准差的数值就更加稀少了。

图4-32显示了正态分布，正态分布是如此的常见，以至于它被认为是自然界的基本特征。它是在相同的情况下，通过对很多独立的因素添加额外因素的方式形成的。例如，心脏的收缩血压受到遗传因子、日常饮食、体重、有氧条件以及其他大量因素的影响。当这些因素开始相互作用的时候，随机挑选的受试者的血压概率分布将会呈钟形分布。

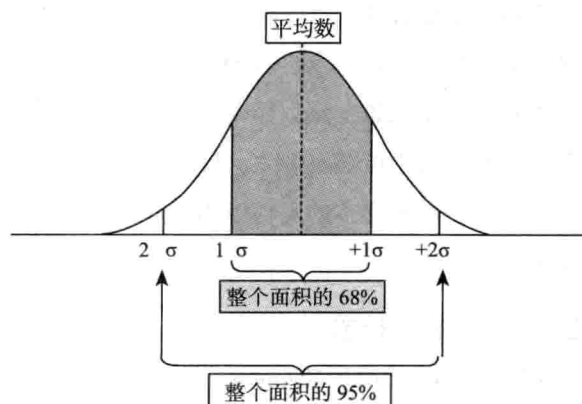


图4-32 正态分布

图4-33³⁹显示了随着平均数抽样分布的形态所出现的各种非正态总体分布情况。请注意，不论在总体中的数据的分布形态如何，当样本规模增加时，抽样分布都会向正态分布集中。数字30在这里并没有任何神奇之处。被视为样本规模函数的从抽样分布向正态分布汇集的比率，是以总体的形态为基础的。注意图中前3个案例，在样本规模为5的水平下，它可以产生几近正态分布的平均数抽样分布。然而，最下面的那个案例，其平均数抽样分布并不能够形成正常形态的抽样分布，只有在样本规模达到30的水平时，才会产生标准的平均数抽样分布。

平均数的标准误差

根据在经典统计学中的表现，我们现在可以进入到确定平均数的抽样分布的最后一步了。到目前为止，我们已经确定的内容如下。

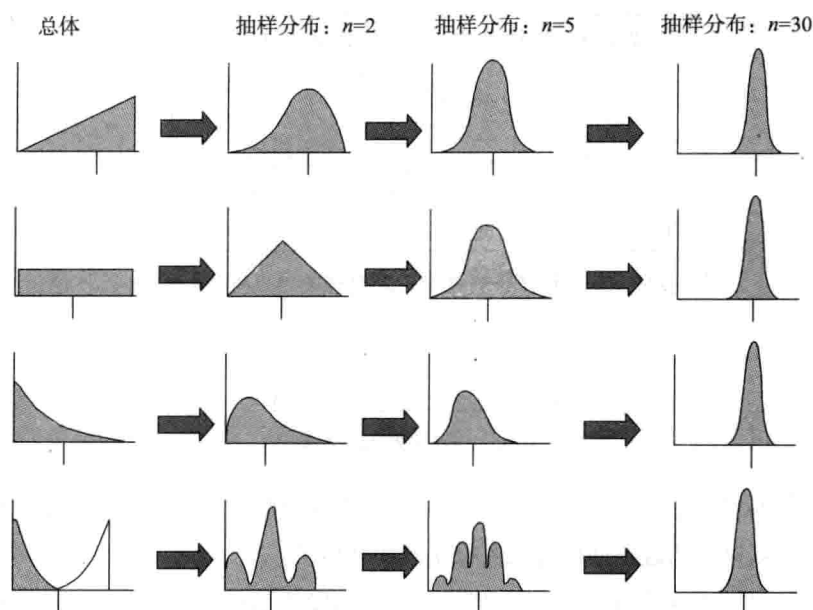


图 4-33 通过给出足够的样本规模, 不管变量分布的形态如何, 平均数的抽样分布达到了正常的形态

- (1) 当样本规模增加时, 平均数的抽样分布会向正态分布的形态转变。
- (2) 正态分布由它的平均数以及标准差来描述。
- (3) 平均数抽样分布的标准差, 也被称为平均数的标准误差, 与从样本中提取的总体的标准差成正比。
- (4) 平均数的标准差与样本规模的平方根成反比。

平均数的标准差等于样本规模的平方根除以总体的标准差。这个表述是正确的, 但是其不太可能应用, 因为总体的标准差是未知的。由此我们可以得出这个总体是不能被观察的结论。然而, 总体的标准差是可以从样本中推算出来的。我们把它称作 Σ 或者 σ , 该公式在图 4-34 中显示。 Σ 或者 σ 上面的帽子状符号对于已经计算出来的数量来说, 是一个统计符号。你会注意到根号里的除数是 $n-1$, 而不是 n 。这种改进弥补了总体的标准差夸大样本的标准差的事实。

这个表达式用来评估总体的标准差, 然后再将其用于图 4-35 中来评估平均数的标准误差。

基于观察资料的样本，对总体的标准差进行评估

$$\hat{\sigma} = \sqrt{\frac{\sum (x_i - \bar{x})^2}{n-1}}$$

其中： x_i 代表在总体的变量 x 中的独立观察资料； $\bar{x}$ 是变量 x 的样本平均数； n 是样本中观察资料的数量

图 4-34 基于观察资料样本，对总体的标准差进行评估

$$\text{平均数的标准误差} = \frac{\hat{\sigma}}{\sqrt{n}}$$

图 4-35 平均数的标准误差

我们已经得到了平均数的标准误差，然后假设抽样分布处于常态，现在就可以对平均数的抽样分布进行正确的评估了。

传统的方法存在着一个问题。它假设抽样分布的形态是正常形态（钟形曲线）。如果情况不是这样的话，那么计算的结果就可能是不精确的了。出于这个原因，我在计算机模拟的基础上，选择一种可以替代的方法来评估平均数的抽样分布。

用计算机模拟计算的方法获得抽样分布

30 年前，我们仍然还没有找到能够替代从单一的样本中获取抽样分布的方法。现在，我们终于找到了一种替代方法，即用计算机模拟计算的方法。我们通过有计划地重新抽样可以多次获得单一的样本，这个结果很有可能就是抽样分布的形态。我们在接下来的部分将讨论这个过程。

我们在第 5 章会提出两种基于计算机的方法：自助再抽样和蒙特卡罗置换。尽管每一种方法都有各自的局限性，但是对于任何一种统计学方法来说它都具有真实性，所以这两种方法可以作为经典方法的替代性方法。

自助再抽样之所以被这样命名，是因为它可以通过辅助程序对自身进行证明，进而更加接近抽样分布的形态。在本书中，这种方法将用来对没有预测力的法则的平均收益率的抽样分布形态做出估计。蒙特卡罗置换是以欧洲最著名的赌场蒙特卡罗命名的，因为这座赌场对计算机的使用与对

轮盘赌的操作一样熟练。

蒙特卡罗置换方法所产生的任何意义，都是以服务抽样分布为目的的。它告诉我们在随机抽样的影响下，没有预测力的法则在回溯测试中产生的收益率大于 0 或者小于 0 的程度。因此，这种方法可以当作绩效的衡量基准。如果某一法则的平均收益率因为抽样变异性而导致太高的话，那我们就可以判断该法则具有预测力。

关于下一章

我们已经引用大量的参考并运用统计学推理的方法对提出的观点进行测试。我们把这些行为用更加正式的术语“假设检验”来定义。在下一章的内容里，我们将讨论“假设检验”的基本原理及其运作方式。

我们还要分析统计学推理的另外一种用法，即置信区间。置信区间的范围包括在一定置信度水平下，总体参数的真实数值。例如，如果被观察的法则的平均收益率为 15%，那么置信区间水平为 95% 就说明该法则的实际收益率在 5% ~ 25%。我们从同样的抽样分布中获得的置信区间的宽度是 20%，那么该数据就可以用于假设检验。

第5章

Chapter 5

假设检验和置信区间

统计学推理的两种类型

统计学推理包括两个步骤：假设检验（hypothesis testing）和参数估计（parameter estimation）。这两个步骤都涉及总体参数的未知值，当一个数据样本与总体参数相关的数值一致或者矛盾时，也就是这个值等于或者小于零时，则该假设检验就会确定。而另外一个步骤——参数估计，则利用样本信息来确定总体参数¹的大概数值。因此，假设检验告诉我们效应是否会显现，而评估则告诉我们效应量。

在某种程度上，这两种推理的形式类似，它们都试图从提取的样本中得出整个总体的结论。除我们已经知道的内容，假设检验和参数估计都在从样本统计量的确定值到总体参数的不确定值的过程中实现了归纳的飞跃（inductive leap）。

当然，参数估计与假设检验之间也存在着差异，因为它们各自的目的不同。对假设检验来说，它的目的在于对接受或者拒绝接受该种推测的总体参数的真实性做出评估；参数估计则致力于提供总体参数的合理数值或数值的范围。从这个意义上讲，参数估计就是一个大胆尝试，并且为我们提供了更多有价值的信息。参数估计不仅仅是告诉我们是否应该接受或者拒绝某种具体的观点，如一个技术分析法则的平均收益率小于零或者等于零，而且还为我们估算某种技术分析法则的平均收益率，并提供一组数值范围，以及该法则的真实平均收益率应该落在哪个特定的概率水平。例如，某法则的预期收益率是10%，那么其收益率落在5%～15%这一范围

的概率为 95%。上述内容中的评估有两种类型：点估计值（point estimate），即该法则的收益率为 10%；区间估计值（interval estimate），即该法则的收益率在 5% ~ 15%。我们在第二部分里讨论的规则就是以估计的方法作为假设检验的辅助工具的。

假设检验 VS. 非正式推理

如果某法则在历史样本中一直都保持盈利，那么这个样本统计量就是不容置疑的事实。然而，在这个事实当中，我们可以对该法则未来的绩效做出推断吗？到底是因为该法则具有真正的预测力而使得它能够获得盈利，还是因为其过去的盈利只是出于运气的成分呢？假设检验对于确定哪种选择才是正确的来说，是一个正式且严谨的推理过程，它有助于回答我们在实际交易中对运用该法则是否合理的疑问。

确认性实证：合适、必需，还不够

在第 2 章中我们曾经指出，因为非正式的推理存在偏差，因此它被确认性实证所替代。也就是说，当我们用常识来对一种观念的有效性进行检验的时候，我们更倾向于寻找确认性实证，即事实与这种观念的真实性相一致。同时，我们也忽视了，或者是很少关注矛盾性实证。常识告诉我们，如果某种观念是正确的，那么这种观念所产生的影响（确认性实证）就是理所当然的。然而，非正式的推理却做出了错误的假设，即确认性实证足以证明其自身就是真理。这是一种逻辑错误。确认性实证不能强行做出某种观念就是真实正确的结论。因为确认性实证只是与这种观念的事实相一致，所以它表现的仅是这种观念的真实性的概率。

必要性实证与充足的实证之间存在着重要的差异。我们曾经使用下面的例子在第 4 章中进行过讨论。假设我们要检验一个论断的真实性：**我看到的这种动物是狗**。我们观察到这种动物有四条腿（实证）。这个实证与（确认）这种动物是狗的表述相一致。换句话说，如果这种动物是一条狗，那么它肯定有四条腿。然而，“四条腿”并不能作为充足的实证来证明这种

动物就是狗，它很可能是其他长有四条腿的动物，如猫、犀牛等。

时下，技术分析领域流行的观点是，试图通过提供某种形态可以成功地预测来证明该形态具有预测力。可以这样说，如果某种形态具有预测力，那么就肯定会有该形态做出成功预测的历史案例。然而，即使我们掌握了这些确认性实证，但从逻辑上来讲，也并不是以证明这种形态就真的具有预测力。我们既不能强行得出一种形态具有预测力的结论，也不能仅凭有四条腿就得出这种动物是狗的结论。

确认性实证能够充分证明确认结果的谬误。

如果 p 是真实的，那么 q 也是真实的。

q 是真实的。

无效的结果：因此， p 是真实的。

如果这种形态具有预测力，那么过去肯定存在成功的案例。

过去确实存在成功的案例。

因此，这种形态具有预测力。

第3章内容表明，尽管确认性实证并不足以证明某一论断的真实性，但矛盾性实证，即与论断的真实性相悖的实证，却足以证明这个诊断是错误的。没有四条腿的动物这个事实足以排斥这种动物是狗的论断。我们把这种有效的辩论形式称为对结果的否定。

如果 p 是真实的，那么 q 也是真实的。

q 是不真实的。

因此， p 是不真实的（错误的）。

如果这种动物是狗，那么这种动物就会有四条腿。

这种动物没有四条腿。

因此，这种动物就不是狗。

假设检验的逻辑偏差是对结果的证伪。因此，这种方法就成为了排斥非正式推理的确认性偏差和有效避免错误观念的有力手段。

什么是统计假设

统计假设就是对总体参数的数值进行推测：通常情况下是数字，如某一法则的平均收益率。总体参数的数值是未知的，这是因为它是不可观察的。由于我们之前讨论过，因此假设它的值等于或者小于零。

观察者知道从总体中提取的样本值就是样本统计量的值，因此，观察者就面临着这样一个问题：被观察的样本统计量的值是否与总体参数的假设值相一致呢？一方面，如果被观察的值靠近假设的数值，我们就应该得出假设是正确的这一合理推论。另一方面，如果样本的数值与假设的数值相差悬殊，那么我们就对假设的真实性提出质疑了。

靠近与远离是两个模糊的术语。假设检验对二者进行了量化，并使得它们可以对假设的真实性做出推断。我们将检验的结果设定在 0 与 1.0 之间，这组数字显示了在假设的数值（已知或者假设）是真实的情况下，被观察的样本统计量偶然发生的概率。例如，假设某法则的预期收益率等于 0，但是在回溯测试中却产生了 +20% 的收益率。假设检验的结果说明有可能会发生以下的事情：如果这个法则的预期收益率确实等于零，那么在回溯测试中，由于存在偶然因素，出现等于或者大于 +20% 收益率的概率就是 0.03，而出现 20% 收益率的概率只有 3%。如果该法则缺少预测力的话，那么我们会信心十足地说这个法则在回溯测试中的表现并不是简单地依靠运气。

用荒谬的实证歪曲假设

假设检验是以正在被测试的假设是真实有效的这一假设开始的。在这个假设的基础之上，预测是根据对各种新的观察资料出现的可能性的假设而推断出来的。换言之，如果一个假设是真实的，那么当其他的结果不太可能出现的时候，某些特定的结果就有可能出现。在这种预期的指引之下，观察者可能利用后续的观察来与预测进行比较。如果预测的结果和观察的一致，那么就不需要对这种假设有所怀疑。然而，如果观察的结果出现了低概率的情况，即结果与假设的真实性不一致，那么这个假设就肯定是错误的。所以，意想不到的实证的出现是排斥一种假设的基础之所在。尽管

这种推理方法有些反常，但从逻辑上来说它是正确的（对结果的否认），也是非常有力的。我们称它为科学发现的逻辑基础。

我们来看一个具体的例子。假设我认为自己是一名非常出色的业余网球选手，于是，我加入了一家网球俱乐部，这个俱乐部会员的年龄和球龄与我差不多。那么基于我的假设，我可以很自信地预测，当我与其他俱乐部选手进行比赛时，我最少能够赢得 $3/4$ 的比赛（胜率等于 75%）。这个预测仅仅来自我的假设。为了检验我的假设，我对自己的前 20 场比赛进行了跟踪。在这 20 场比赛过后，我感到既震惊又失望。我不仅没有赢得一场比赛的胜利（胜率为 0），而且大部分比赛都处于绝对劣势。这一结果与我从自己的假设当中所推断出来的结果大相径庭。也就是说，我所做的假设暗示了出现我所预期的概率是非常低的。这种出人意料的实证的出现迫使我以前假设进行修正（证伪），并且让我就此放弃了对网球事业的美好幻想。²

在上述情况下，实证是绝对清楚的，即我输掉了全部 20 场比赛。然而，如果实证表现的不太清楚，又会是什么样子呢？假设我赢得了全部比赛的 $2/3$ ，即我的胜率为 0.66，略低于预测的 0.75 的胜率。到底是因为与预测的胜率相背离而随机出现负值，还是由于存在巨大的离差致使我自己打网球能力的假设是错误的呢？统计分析在这里就显得非常重要。统计分析试图回答这些问题：是因为被观察的胜率（0.66）与预测的胜率（0.75）之间的差异太大，而增加了对假设的真实性的质疑程度；还是因为在网球比赛这种特殊的样本中，0.66 与 0.75 之间的差异仅仅是随机变量的缘故呢？假设检验试图区分两种预测错误，即因为离差太小而导致出现随机抽样的结果，以及因为离差太大而导致错误假设的出现。

双重假设：原假设 VS. 备择假设

假设检验依赖于间接证明的方法。也就是说，它通过证明其他事物是错误的来证明某种事物的真实性。因此，为了证明我们想要证明的假设是正确的，我们就要证明与其相反的假设是错误的。为了证明假设 A 是正确的，我们就要证明相反的假设 Not-A 是错误的。

假设检验包括两种假设：一种叫作原假设，而另外一种叫作备择假设。虽然这两种假设的名字听上去有些奇怪，但它们却广为人知，故为本书所用。科学家们运用备择假设来证明重要的新发现，而原假设则只证明没有发现新的知识。例如，乔纳斯·索尔克（Jonas Salk），脊髓灰质炎疫苗的发明者，提出了备择假设，他新发明的疫苗在预防脊髓灰质炎上要比安慰剂更加有效。原假设则主张索尔克的新疫苗与传统的安慰剂相比，并不能有效预防脊髓灰质炎。对于本书提到的技术分析法则而言，备择假设认为预期收益率要大于零，而原假设则主张预期收益率小于零。

为方便起见，我将采用传统的符号： H_A 代表备择假设， H_0 代表原假设。因为原假设主张没有发现新知识，所以就用 H_0 来表示。

假设检验逻辑的关键是对 H_A 和 H_0 的定义，它们之间虽然相互排斥，但却覆盖了所有命题。后者是什么意思呢？如果把它们两个结合在一起，就会覆盖所有的可能性。不论脊髓灰质炎疫苗是否能够起到预防的作用，肯定不会出现第三种可能。就是说，在技术分析领域，一种技术分析的收益率只有大于零和小于零两种结果。

两种假设又是相互排斥的。相互排斥的命题不可能在同一时间被证明都是正确的，如果 H_0 被证明是错误的，那么 H_A 就肯定是正确的，反之亦然。通过对两种假设既相互排斥又覆盖全面的这一描述，那么当其中的一种假设被证明是错误的的时候，我们就可以得出另外一种假设是正确的这一肯定的结论。我们把这样的证明称为间接证明。这些概念如图 5-1 所示。

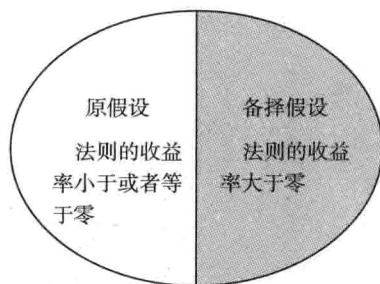


图 5-1 相互排斥而又彻底的假设

假设检验的基本原理

我们有必要对检验假设的两个方面进行解释。首先，为什么检验主要集中在原假设身上？其次，为什么要假设原假设比备择假设更加正确？以下内容将会对此进行解释。

为什么原假设会成为检验的目标

我们在第3章中已经讨论过, 实证可以从逻辑上来推断一种假设是错误的, 但是它不能推断该假设是正确的,³ 因此, 假设检验必须证明有些东西是错误的。但问题是: 在 H_A 和 H_0 两种假设当中, 到底哪一种假设会成为实际的目标呢?

在两种矛盾的观点中, H_0 成为证伪的目标。这是因为, H_0 可以归纳有关参数值的单一论断, 这也意味着只执行一个检验。如果这个单一值可以成功地对实证发起挑战, 那么 H_0 就是错误的。相反, H_A 代表了对参数值进行论断的无限数量, 因为并不针对特定数值, 所以检验的无限数量将会对备择假设进行证伪。

实际上, H_0 和 H_A 代表了法则预期收益率的无穷多论断, 但是 H_0 可以归纳为单一的论断。第一, 我们来分析一下为什么 H_A 代表了无穷多的论断。在法则的收益率大于零的表述中, H_A 有效地说明了法则的预期收益率, 而在即将实现预期的过程中, 它可以是任何一组无限大于零的数值: +0.1%, +2%, +6%, 以及其他为正的数值 (见图 5-2)。假设 H_A 产生了无穷多的论断, 那么无限数量的检验就会令人信服地对其进行排斥。很明显, 这种做法是不切实际的。

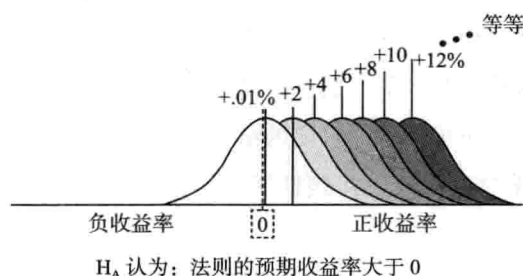


图 5-2 备择假设对法则的预期收益率做出了无穷多的论断

H_0 也会产生无穷多有关总体参数值的论断, 它认为法则的平均收益率等于零或者小于零, 其中法则的平均收益率等于零的论断最重要。对单一的论断进行证伪才是实际的目标。如果这些论断 (收益率等于零) 的大多数正值可以通过法则的回溯测试概率加以排斥, 那么所有小于论断

的值（如法则的收益率为 -1.3% ）就是矛盾的，并且它们会向更深的程度发展。这就是 H_0 可以归纳出法则的收益率为零这一单一论断的原因（见图5-3）。

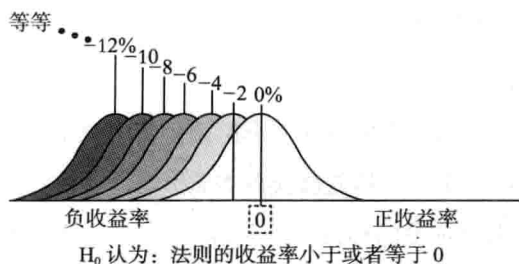


图 5-3 原假设对总体平均值做出了无穷多的论断,但是只有一个结果

将原假设视为正确的原因

假设检验将 H_0 视为正确的原因基于以下两个方面：科学的怀疑态度和简单原则（简约性）。

1. 怀疑的态度

对 H_0 是正确的假设与对所有新知的观点抱有科学的怀疑态度是一致的，我们已在第3章中就此进行过解释。我们之所以采取这种保守的立场，是因为表达某种观点很容易，但是要做到对其进行真正的挖掘则是相当困难的。对其进行证明的压力势必落在提出这种观点的人的身上。

科学对知识宝库免受谎言和奇怪信仰的侵害，表现出了合理的担忧。我们可以在刑事诉讼体系中找到类似的案例。自由社会对防止个人遭受国家大范围的检控表现出了合理的担忧。出于这种原因，刑事诉讼中被告的首要任务就是做出无罪推定，这就使得证明其有罪的压力大增，即证明无罪推定的假设是错误的。在陈述过程中，这种压力无疑是非常大的。为了让被告获得有罪判决，陈述必须在合理的怀疑之外提供犯罪的证据，这个要求确实太高了。初始的无罪推定假设和提供佐证的高要求，使其成为无辜者避免因为遭受侵害而面临死亡和牢狱之灾的法制途径。

按照科学传统，假设检验对于新知的证明面临着巨大压力。实际上，当某一法则在回溯测试中的绩效大于零时，样本的实证则偏向 H_A ，即法则

的预期收益率大于零。然而，假设检验需要的不仅仅是合理的证据，它更需要在放弃对 H_0 做出真实的假设之前，获得引人注目的实证。关于技术分析法则具有预测力的主张 (H_A)，在其未被从事技术分析的科学实践者接受时，对其进行证明会面临巨大压力。

2. 简单原则

对原假设给予优先考虑的另外一个原因是简单原则。这一基础的科学原则认为，理论越简单，就越能够获得本质形态，这一点是缜密的理论所达不到的，我们把这种原则称为“奥卡姆剃刀”。也就是说，如果一种现象能够被不止一种假设所解释，那么其中最简单的那种假设就很有可能是正确的。例如， H_0 对某法则过去获得的成功解释为运气的成分，而 H_A 则认为获得的盈利来自具有预测力的重复出现的市场形态，相比之下， H_0 比 H_A 的解释更简单。

愈加简单的解释（理论、法则、假设、模型等），愈有可能成为正确的事实，这是因为它们不可能是偶然的数据拟合。越复杂的解释，即大量的假设、条件以及局限性，就越有可能是一组观察资料的数据拟合。这表明，当把一个数学函数代入一组数据时，一个函数可以被视为对一组给定的观察资料进行解释的数学假设。如图 5-4 所示，两个数学函数被代入同一组数据当中，其中一个函数是线性的。之所以说它简单，是因为它由两种系数或者自由度定义的：倾斜度和它与纵轴（Y 轴的截距）的切点。换句话说，就是与数据相匹配的线可以通过对这两个因素的控制得到证明。

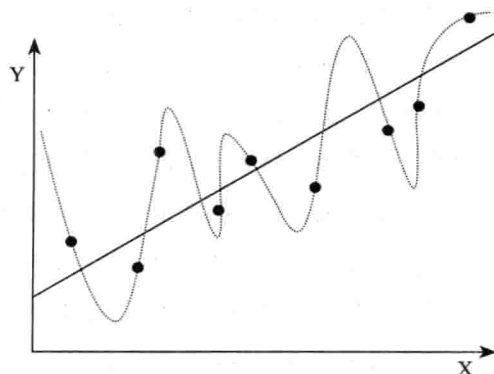


图 5-4 简单的优越性：所有其他因素都一样

另外一种函数就是带有 10 个系数的多项式函数。每一个系数都能够让曲线发生弯曲。随着弯曲程度的不同,曲线可以迂回穿行,并且触碰每一个数据点。实际上,当运用“最小二乘回归”的方法把数据代入函数时,如果函数可以包含和自由度(系数)同样多的数据点,那么完美的匹配就会得到保证。尽管线性函数不能触碰到每一个观察点,但是它能够描绘出数据的总趋势,我们把正在增加的 X 值与 Y 值联系起来。换句话说,简单函数可以捕捉到数据的本质特征,并且可以描述 X 和 Y 之间的实际关系。相反,复杂函数最有可能对数据的随机波动做出详细描述,并显示 X 和 Y 之间的正相关关系。

当然也不能说复杂的曲线就是永远不能修正的。如果有足够的数据,曲线就会通过复杂的过程而产生,那么一个缜密的模型就会得到很好的修正。然而,在其他条件保持不变的情况下,越简单的解释就越有可能是正确的。

有力和无力决策

假设检验导致只有一种决策:拒绝 H_0 或者保留 H_0 。这两种决策存在着本质上的区别:第一种是有力决策,而第二种则是无力决策。⁴ 拒绝 H_0 属于有力决策,它受到了与 H_0 发生矛盾的概率小,并且信息丰富的实证的影响。相反,保留 H_0 的决策则是相对次级的选择,因为实证不过是与我们的预期和假设的内容相符合而已,即 H_0 是正确的。换句话说,检验统计量的值并不出人意料,也没有包含多少信息。

倘若索尔克的脊髓灰质炎疫苗的受试者发现其与安慰剂的使用者将会面临同样的患病风险时,那么 H_0 就是正确的,因此,我们对这种疫苗的无效就不会感到意外了。如果这一切都没有发生,我们无疑是在实验室中又度过了令人失望的一天。当然,这一切很难真的发生,要不医学历史的进程就会因此而改变了。索尔克做出拒绝 H_0 的决定是受到接受治疗群体的感染率的影响,而接受新疫苗治疗的人数要比接受安慰剂治疗的人数少得多。

由此看来,相对于做出保留 H_0 的决策,拒绝 H_0 从根本上来讲更加有力。我们之所以做出这一决策,是因为手中的实证十分充分,可以对新知识的

质疑予以排斥。相反，保留 H_0 的决策是由于缺少有力的实证而做出的。缺少有力的实证并不意味着原假设就一定正确，或大致正确，⁵ 它只意味着原假设有可能正确。由于科学对于新知识持有保守的立场，在缺少有力实证的情况下所得出的结论，我们可以将其理解为没有新的发现。

假设检验：构成法

假设检验的三个重要组成部分是：①假设；②检验统计量；③在假设是真实的条件下，获得产生检验统计量大概分布（抽样分布）的手段。⁶ 术语“检验统计量”是指对假设进行测试的样本统计量。因此，这些术语可以相互替代。

让我们回顾一下有关法则测试的内容：①原假设法则的预期收益率等于或者小于零；②检验统计量就是在历史数据样本中对法则进行回溯测试时获得的平均收益率；③抽样分布代表着，如果当法则被用于许多独立样本的测试时，该法则平均收益率的随机变量。抽样分布正中间的数字 0 代表平均收益率为 0，这反映了原假设的观点。

我们曾经在第 4 章中讨论过，抽样分布可以通过两种途径获得：经典的统计学分析方法和计算机模拟。以计算机模拟为基础的方法有两种：蒙特卡罗置换和自举法。这两种方法都会的第二部分的案例研究中得到应用。

疑问：检验统计量是不太可能出现的吗

假设检验的基础思想很简单：在某一假设是真实的情况下得出的罕见的结果（观察资料），就是证明该假设不是真实的有力实证。⁷ 如果我的假设“我是一名出色的网球选手”是正确的，那么对于我连续输掉 20 场比赛的事实来说，这个结果就是非常罕见的（不太可能出现的）。令我吃惊和难堪的是，我确实连续输掉了 20 场比赛，这一实证预示着我的假设就是错误的。

假设对一种名为“移动平均 50”的法则进行回溯测试。我们对“移动平均 50”法则的定义是这样的：如果标准普尔 500 指数的收盘价大于 50

日移动平均值,那么就持有一个标准普尔 500 指数的多头头寸;反之,就持有一个空头头寸。备择假设认为,该法则具有预测力,且预期收益率在消除数据趋势的情况下大于 0。⁸ 0 是指没有预测力的法则的预期收益率。原假设认为,“移动平均 50”法则的预期收益率等于或者小于 0。原假设的观点如图 5-5 所示。横轴代表即将实现的预期收益率。

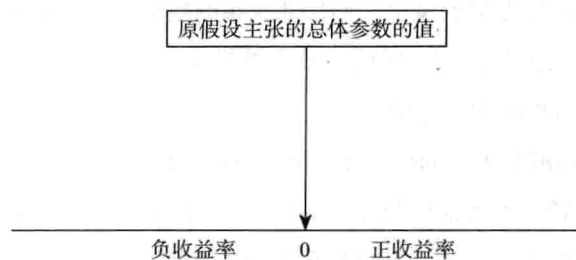


图 5-5 原假设主张的总体参数的值

“移动平均 50”法则的平均收益率是从对历史数据样本中的法则进行回溯测试获得的。当把这个数值作为总体参数的假设值绘制在同样的横轴上时,我们就得到了图 5-6。

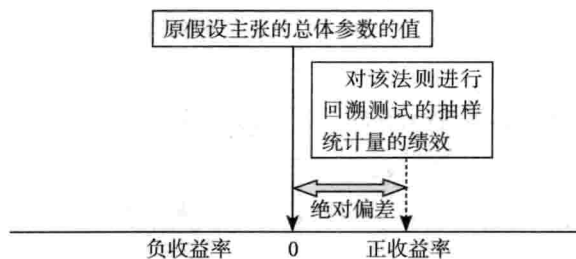


图 5-6 与回溯测试的绩效进行比较的总体参数的假设值

注意图 5-6 中,对于由原假设预测的值与从回溯测试中获得的绩效之间的正绝对偏差值,我们提出了这样一个疑问:出乎我们意料的正绝对偏差值,是不是应该作为不合理的假设被拒绝呢?

对于正绝对偏差值,我们可以做出如下两种解释:第一,正绝对偏差值是由抽样误差造成的,即法则在使用回溯测试特定的数据抽样时,存在很大的运气成分;第二,或者是由于 0 的假设值是错误的,即该法则确实具有预测力,且它的预期收益率也的确大于 0。如果当实证,特别是正绝

对偏差值的容量是非常罕见且出人意料,或者是不太可能出现,进而可以排斥 H_0 的时候,假设检验的目的就达成了。

为了评估实证的不可能性,被观察的平均收益率,特别是它与假设值的背离情况,是通过抽样分布的方式来进行评估的。回想一下,抽样提供了在被观察的抽样统计量的值与由于抽样误差得出的预期收益率之间出现各种容量的绝对偏差值的概率。如果被观察的数值的离差大于因为抽样误差而导致的数值的话,那么原假设就会被排斥,而备择假设就会被接受。换言之就是,这种法则具有预测力。

图 5-7 显示的是由于抽样变异性而导致出现的样本统计量的值。注意,抽样分布的位置很有可能使收益率的值为 0,这仅仅反映的是原假设所主张的总体参数的值。同时,我们还要注意落入抽样分布随机变量范围之内的被观察(回溯测试)的平均收益率。这个结果并不出人意料。换句话说,在被观察的检验统计量的值与假设的(预测的)值之间出现的离差,有可能就是由于抽样误差导致的。因此,并不充分的实证是无法排斥原假设的。

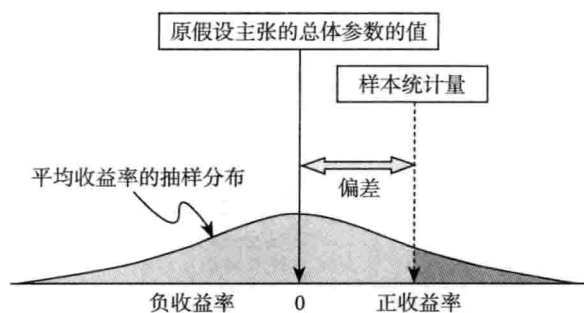


图 5-7 出人意料的实证位于抽样变异范围内,原假设没有被排斥

如果被观察的收益率与假设的收益率之间的离差太大,从而导致对原假设予以排斥时,抽样分布的宽度在此时就起到了决定性作用(见图 5-7)。我们在第 4 章中看到,统计量抽样分布的宽度是由两种因素决定的:①提供抽样整体范围内的变量总和;②包含抽样在内的观察资料的数量。关于第一种因素,包含总体数据的变异性越大,每日收益率抽样分布的宽度也就越大。至于第二种因素,包含抽样的观察资料的数量越大,抽样分布的宽度也就越小。

在图 5-8 中, 抽样分布相对狭窄。被观察的样本统计量的值位于抽样分布右侧尾部的外面。我们可以认为这个观测数据是不太可能发生的, 或者是让人感到意外的, 而且这个数据与假设的数值不一致, 所以这个实证就可以排斥原假设。

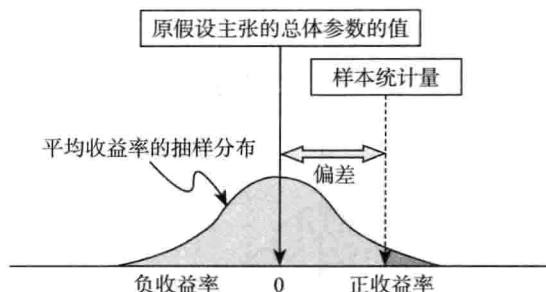


图 5-8 位于抽样变异范围以外的让人感到意外 (不太可能发生的) 的实证, 原假设被排斥

上面的几幅图表对被观察的数值与假设的数值 (假设预测的数值) 之间离差的容量是如何对假设进行证伪做出了直观的总结。更加严谨地说, 这种直觉必须要被量化。这就需要把被观察数值的离差转变为概率, 特别是在假设的数值是真实的条件下, 给出大的离差值的概率。这一概率要视特殊条件而定, 即假设的数值是真实的, 我们把它称作条件概率 (conditional probability)。换一种说法就是, 条件概率是建立在其他事实是真实的这一基本条件之上的。

在假设检验中, 条件概率有一个特殊的名字, 即概率值 (p-value)。具体来说就是, 观察到的检验统计量的值在被测试的假设是正确的基础上可能发生的概率。概率值越小, 对原假设的质疑也就越合理。如果概率值小于某一临界值, 而这个值又是在检验结果出来之前必须确定的, 那么原假设就会被排斥, 而备择假设则会被接受。我们也可以将概率值理解为, 当原假设确实真实有效时, 其被错误排斥的概率。概率值的另外一种表示方法是图释。概率值等于抽样分布位于大于且等于观察到的检验统计量的值的全部面积。

我们现在来分析一下与法则测试有关的内容。如果一种法则在回溯测试中的收益率为 +3.5%, 我们把这个值绘制在代表抽样分布的横轴上, 然

后再确定抽样分布的面积占据等于或者大于 3.5% 的部分。假设这个面积等于抽样分布全部面积的 0.10, 0.10 就是抽样统计量的概率值。这一事实等于是说, 如果法则的收益率为 0, 那么该法则在回溯测试中由于抽样变异性(偶然)而导致等于或者大于 3.5%, 这一数值的概率就是 0.10, 具体情况如图 5-9 所示。

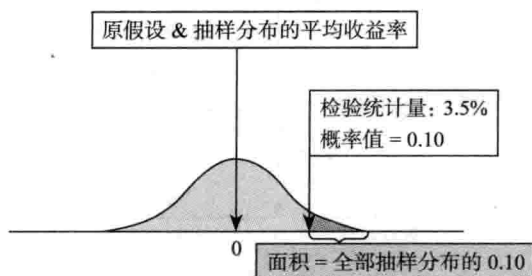


图 5-9 概率值: 抽样分布的面积大于 +3.5%, 3.5%, 或者比 3.5% 更高的条件概率, 说明原假设是正确的

概率值、统计显著性以及对原假设的排斥

检验统计量的概率值的另一种叫法是“检验的统计显著性”。概率值越小, 检验结果的统计显著性就越明显。统计显著性所导致的低概率值完全可以对排斥原假设做出解释。

检验统计量的概率值越小, 我们就越有信心拒绝原假设是正确的决策这一观点。概率值可以表示检验统计量的值符合原假设的程度。概率值越大, 符合的程度就越高, 反之亦然。换种说法就是, 即相对于给定的观点(假设), 观测数据越出人意料(不太可能发生), 这种观点就越有可能是错误的。

到底概率值要多小才能够解释对原假设的排斥呢? 这不仅是一个具体的问题, 而且还会涉及由于错误排斥而导致付出的代价。以下部分就涉及这方面的内容。然而, 还有一些被普遍使用的标准需要我們注意: 概率值为 0.10, 通常称为可能显著; 0.05 或者更小的概率值则称为统计显著性, 并且被视为最大的概率值, 它们可以对原假设进行排斥; 当概率值为 0.01 或者更小时, 称为非常显著; 当概率值为 0.001 或者更小时, 称为高度显著。

检验结论和错误

假设检验可以导致两种结论当中的任何一种：①排斥原假设；②接受原假设。如果假设是真实的，也会出现两种可能：①原假设是真实的，也就是说，法则的预期收益率等于或者小于0；②原假设是错误的，即法则的预期收益率大于0，这是因为它拥有一定的预测力。考虑到检验及真实性都会出现两种可能的结论，因此，总共就会出现4种可能的结论。图5-10中的4个单元格里显示的就是上述内容。

		真实的 (实际上只有上帝才知道)	
		原假设是正确的 法则的收益率 ≤ 0	原假设是错误的 法则的收益率 > 0
测试结果与结论	预测值高，原假设没有被排斥	正确的结论 技术分析法则无效，放弃它	第二种类型的错误 技术分析法则有效，但我们没有使用，导致错失了机会
	预测值高，原假设被排斥	第一种类型的错误 技术分析法则无效，经过使用，获得的收益率为0，还得承受一定的风险	正确的结论 技术分析法则有效，通过使用而获利

图 5-10 假设检验可能出现的 4 种结果

假设检验也会导致两种类型的错误：第一种类型的错误是指，当概率值较低的时候，我们认为原假设是错误的，而实际上，原假设是正确的。也就是说，当法则确实缺少预测力时，由于运气的成分导致其在回溯测试中产生了丰厚的利润，那么出现这种较低的概率值的时候，我们就有理由对原假设做出排斥。这种错误的出现，完全是由于法则的研究者受到了随机性的欺骗。第二种类型的错误是指，当较高的概率值让我们将原假设保留下来时，而实际结果却是完全错误的。也就是说，回溯测试误导我们得出了法则不具有预测力的结论，而实际上它的预期收益率却大于0。

当假设检验正在进行的时候，只有上帝才知道会发生哪种类型的错误。而更多依靠统计推断的凡人则必须接受这样的现实，即检验的结果也许会出现错。

站在技术人员的客观立场上，两种类型的错误会导致出现不同的结果。在第一种类型的错误中，原假设被错误地排斥掉了，从而导致我们运用无效法则，这揭露了营运资本正在冒着没有补偿预期的风险。第二种类型的错误则导致有效法则被忽视，从而失去了交易机会。在这两种类型的错误当中，第一种类型的错误更加严重，失去营运资本要比失去交易机会的性质严重得多，尽管还会有其他的交易机会，但是当资本被耗之殆尽时，自己也会被市场淘汰。

假设检验可以通过两种途径证明其是正确的：当法则没有价值的时候，正确地排斥原假设；当技术分析法则确实无效时，正确地接受原假设。

用计算机模拟的方法产生抽样分布

我们前面曾经提到过，检验假设需要一种能够判断检验统计量的抽样分布的方法，目前有两种方法可以解决这个问题：传统的数理统计学方法，以及最近才发展起来的计算机模拟的随机方法。这部分内容讨论的是两种基于计算机的方法：自举法与蒙特卡罗置换法。

在只有一个单独数据样本的情况下，传统方法和计算机模拟方法都可以解决在检验统计量中判断随机变量的程度这一问题，因此，就只会得出单一的检验统计量的值。正如我们前面所讲的那样，统计量的单一值是不可能表达出变异性的意义的。

计算机模拟方法通过对原始观察资料的随机再采样（重复使用），得到新的计算机生成样本，并据此对抽样分布的形态做出判断。我们要对每一次的再采样计算检验统计量。这一过程可以重复上千次，并且得出一组特大的抽样统计量的值。重复地对原始观察资料进行再采样，可以更加接近抽样统计的变异性。虽然这种方法看上去有点奇怪，但是我们还要运用。因为这种方法不仅可以很好地运用在实际操作中，而且也合理可靠的数学理论打下了基础。

自举法与蒙特卡罗置换这两种计算机模拟方法的共同之处在于，它们都要依靠随机化。也就是说，它们对原始样本进行随机再采样。然而，这

两种方法在很多重要的方面都存在着差异。首先，它们对原假设的不同描述只进行粗略的检验。例如，当原假设认为正在进行检验的法则不具有预测力时，这两种方法对此的检验只是稍有不同。用自举法检验的原假设认为，该法则收益率的总体分布的值小于或者等于0。相反，以蒙特卡罗置换法进行检验的原假设则主张，该法则的输出值（+1 和 -1）是随着未来市场价格的变化而随机搭配的，⁹换句话说，即该法则的输出是无信息的噪声，它可以从轮盘赌中轻易得出。

由于自举法与蒙特卡罗置换法对原假设的不同描述进行检验，因此它们需要不同的数据。自举法利用该法则的每日收益率历史记录；蒙特卡罗置换法则利用该法则的每日输出值历史记录（如 +1 和 -1 的排列顺序），以及交易市场每日价格变化的历史记录。

这两种方法使用的随机抽样方法也不同。自举法使用的是一种被称作返回抽样的随机化方法，而蒙特卡罗置换法则利用不可替代的市场收益率来随机匹配法则的输出值。这一种差别会在我们对这两种方法的运算法则进行描述时谈到。

正是由于这些不同，才使得这两种方法产生的抽样分布也不尽相同。因此，我们通过这两种方法得出的结论有可能会排斥原假设，也可能不会。然而，从蒙特卡罗置换法的开发者蒂莫西·马斯特斯博士（Timothy Masters）所实施的大量模拟来看，这两种方法在消除趋势的市场数据中的表现还是得到了认同的。出于这个原因，本书列举的自举法与蒙特卡罗置换法都使用消除趋势的市场数据。

蒙特卡罗置换法的最后一个特点就是，它为大众公用；而用于法则检验的自举法则受到版权的保护，只能从它的开发者——Quantmetrics¹⁰ 手中获得。

自举法

自举法是由埃夫隆（Efron）¹¹ 在 1979 年首先提出的，随后在埃里克·诺琳（Eric Noreen）的《检验假设的计算机模拟方法》（*Computer Intensive Methods for Testing Hypotheses*）¹² 对此进行了重新定义。自举法通

过可以替代的再抽样方法，从原始样本中找出检验统计量的抽样分布。

自举法是基于真实且让人吃惊的数学事实，即自举法原理。当数学定理还未从已经建立起来的理论和数学体系的基本假设（原理）当中推导出之前还是尚未得到认知的事实。假设在满足了一定合理条件的情况下，当样本容量趋于无限大时，自举法原理可以确保得到一个正确的抽样分布。站在实践的角度来看，则意味着在给定的单一观察资料的样本中，自举法可以产生用于检验技术分析法则的显著性所需的抽样分布。

在这个基本的形式中，自举法并不适合对通过数据挖掘发现的法则进行统计显著性的估算，然而，由加利福尼亚大学圣迭戈分校的经济学教授哈尔伯特·怀特博士（Halbert White）发明并保有专利权的一项针对自举法的修正方法，则将自举法的应用扩展到了通过数据挖掘发现的法则当中。对自举法的这种修正，与软件中所谓的对预报员的真实考验并不一致，这一点我们将在第6章讨论，并在第9章对用于标准普尔500指数的6000多种法则进行统计显著性评估时使用。

自举法的程序：怀特真实检验

我们接下来对用于检验技术分析法则的统计显著性的自举法进行描述，因此，下面的介绍并不涉及数据挖掘的内容。图5-11显示了自举法的程序。每一个再抽样和原始样本之间的双箭头表示抽样是可以被替代的（下面进行详解）。

在图5-11中有几点内容是值得注意的，如下所示。

第一，原始样本用大的椭圆形表示，它包括消除趋势的数据法则的每日收益率，如我们在第1章中讨论的那样，消除趋势的市场数据的平均每日变动为0。

第二，在再抽样操作开始之前，法则的每日收益率会先经过叠合法（zero-centering）调整，请不要将此与消除趋势相混淆。通过叠合法的调整，法则的平均每日收益率为0。换句话说，如果该法则能够在消除趋势的数据中获得结果不是0的收益率，那么它的收益率肯定会向0集中。这有助于让每日收益率与原假设相一致，因为后者认为它们的平均值等于0。这

一步达到了计算法则的平均每日收益率的目的, 然后, 再从法则的每日收益率中减去这一平均值。一旦法则的每日收益率向 0 集中, 就能够产生与原假设相一致的抽样分布。

第三, 包含每一个再抽样的每日观察资料的数量必须准确地等于原始样本中观察资料的数量。如果每一个再抽样中观察资料的数量等于原始样本中观察资料的数量, 我们就可以说自举法原理是正确的。在图 5-11 中, 原始样本包含的观察资料是 1 231 个, 因此, 每一个经过自举法的再抽样也应该包括 1 231 个观察资料。

第四, 每一个再抽样是由可以替代的样本产生的。这意味着, 法则的每日收益率在从原始样本中随机挑选出来之后就已经被记录下来, 然后我们把个值放回原始样本中, 并对法则的每日收益率进行再一次的随机挑选。这就使得在给定的样本中, 单独的每日收益率被选中的可能性要大于那些曾经被选上, 或者从来没有被挑选上的值。因为这种随机因素, 才使得自举法的程序对创建抽样统计量的变异性模型成为了可能。

第五, 图 5-11 显示采用了 5 000 个再抽样。每一个再抽样都计算出一个平均值, 再用这 5 000 个平均值绘制平均值的抽样分布图。

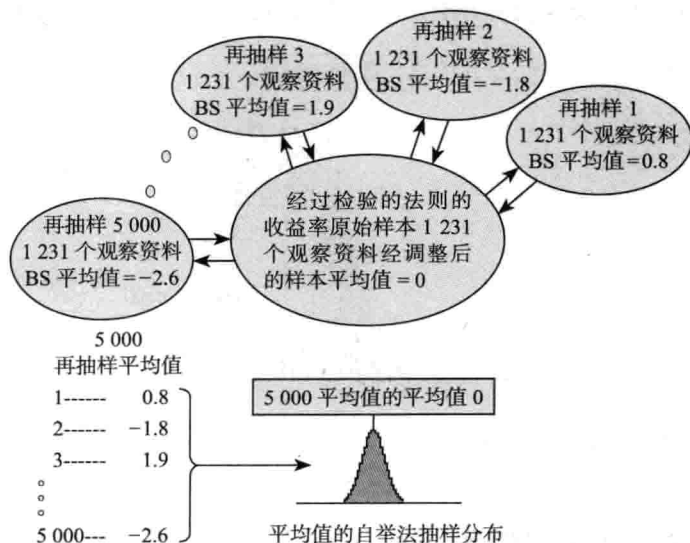


图 5-11 自举法产生样本平均收益率的抽样分布过程

绘制平均值抽样分布图的步骤如下。

(1) 对原始样本(见图 5-11)中 1 231 个观察资料进行法则的平均每日收益率的计算。

(2) 运用叠合法:从原始样本的每日收益率中减去平均每日收益率。

(3) 将 0 集中的数据提取出来单独存放在一个容器里。

(4) 从单独存放的数据中随机挑选一个每日收益率,并记录下它的值。

(5) 把刚挑选出来的值放回去,并将这个容器中的数据重新搅动(有些统计学家更喜欢晃动这个容器)。

(6) 步骤 4 和步骤 5 反复运行 $N-1$ (1 230) 次,并创建总量为 N (1 231) 的随机挑选的观察资料。至此,自举法抽样完成了第一步。

(7) 计算在第一次再抽样中观察资料 N (1 231) 的平均收益率。这个值就是自举法平均值当中的一个。

(8) 运行步骤 6 ~ 步骤 9 共计 5 000 次,可以获得大量的自举法平均值。

(9) 绘制平均值的抽样分布图。

(10) 比较法则的平均收益率和抽样分布,并确定包括抽样分布在内的 5 000 个平均值,这部分要超过决定概率值的平均收益率(见图 5-12)。

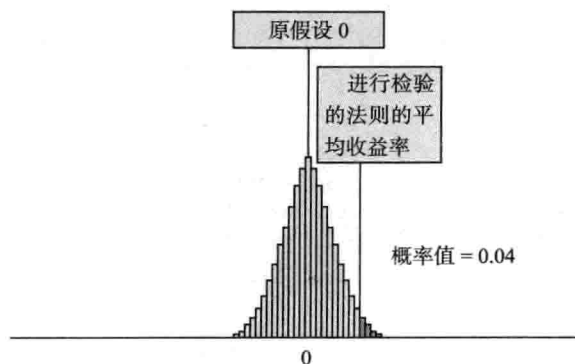


图 5-12 法则的平均收益率与自举法下的平均值的抽样分布的比较

蒙特卡罗置换法 (MCP)

蒙特卡罗置换法是由斯塔尼斯拉夫·乌拉姆 (Stanislaw Ulam, 1909—

1984)发明的, 这种方法是解决随机抽样问题的一般方法。而将蒙特卡罗置换法应用于对技术分析法则的检验则是由蒂莫西·马斯特斯博士自主开发的。他是第一个提出用这种方法产生抽样分布的人, 并检验法则在回溯测试中产生绩效的统计显著性。这种方法被视为“怀特真实检验”的替代法。

自从蒙特卡罗置换法发明以来, 它还从来没有用过对法则的检验, 这就使得马斯特斯博士可以就蒙特卡罗置换法对没有预测力的法则生成抽样分布这一特点开始进行深入的了解。这个任务可以通过对消除趋势的每日市场收益率(如标准普尔 500 指数)与代表每日法则输出值的有序¹³时间序列进行随机配对(序列改变)来完成。回想一下由蒙特卡罗置换法检验的原假设, 该假设认为, 正在进行检验的法则的收益率是从没有预测力的法则所产生的收益率总体中获得的一个抽样。类似这种法则的收益率可以通过对该法则的输出值(+1 和 -1)和市场价格的变化进行随机配对来模拟。像这样的随机配对会使本身可能具有预测力的法则彻底失去预测力, 我把这种随机配对称为噪声规则(noise rule)。

图 5-13 显示了将市场价格的变化与法则输出值进行随机配对的过程。每日法则输出值的时间序列就是那些按照原始序列进行估算的法则所产生出来的时间序列。当输出值的序列与实际上只是市场历史的无序版本进行配对之后, 就可以计算出噪声规则的平均每日收益率。这个数值被放在每一行末尾的灰色方框中。如果要绘制抽样分布, 则需要更多这样的数值。

为了获得更多有关噪声规则平均收益率的值, 可以将实际法则输出值的相同时间序列与大量随机的市场价格变化进行配对(序列改变)。图 5-13 只显示了蒙特卡罗置换法下的 3 种结果, 而在现实中, 这个数量将会更多, 也许是 5 000 个。用这 5 000 个平均收益率的值绘制噪声规则(没有预测力)的平均收益率的抽样分布。

步骤

通过蒙特卡罗置换法产生抽样分布的步骤如下。

(1) 获得技术分析法则进行检验时段的单日市场价格变化的样本。对消除趋势的描述同第 1 章。

法则输出值 的时间序列	+1	+1	+1	+1	-1	-1	-1	-1	-1	-1	
随机得出的 标准普尔 500 指数的收益率 1	-0.8	+0.3	-0.9	-2.6	+3.1	+1.7	-0.8	-2.6	+1.2	-0.4	平均收益率 1 -0.62
随机的法则 收益率 1	-0.8	+0.3	-0.9	-2.6	-3.1	-1.7	+0.8	+2.6	-1.2	+0.4	
随机得出的 标准普尔 500 指数的收益率 2	-0.4	+1.2	-2.6	-0.8	+1.7	+3.1	-2.6	-0.9	+0.3	-0.8	平均收益率 2 -0.34
随机的法则 收益率 2	-0.4	+1.2	-2.6	-0.8	-1.7	-3.1	+2.6	+0.9	-0.3	+0.8	
随机得出的 标准普尔 500 指数的收益率 3	-2.6	+1.7	-0.4	-0.9	-0.8	-0.8	+0.3	+1.2	-2.6	+3.1	平均收益率 3 -0.26
随机的法则 收益率 3	-2.6	+1.7	-0.4	-0.9	+0.8	+0.8	-0.3	-1.2	+2.6	-3.1	

图 5-13 蒙特卡罗置换法

(2) 获得在进行回溯测试期间的每日法则输出值的时间序列。假设经过检验的某一法则的每日输出值是 1 231。

(3) 将前一天市场的价格变化趋势记录在一张纸上，然后把它放入一个容器，并进行搅动。

(4) 从容器中随机挑选出记录市场价格变化的纸张，并将其与第一个（最早的）法则的输出值进行配对。不要将记录市场价格变化的纸张放回容器，也就是说，这个抽样是不需要替代的。

(5) 重复步骤 4，直到容器中所有记录的收益率都与法则的输出值配对完毕。在这个例子中，总计会有 1 231 个这样的配对。

(6) 对随机配对的这 1 231 个组合进行收益率的计算。可以用法则的输出值（+1 代表多头，-1 代表空头）乘以前一天的市场价格变化。

(7) 算出从步骤 6 中获得的收益率的平均收益率。

(8) 重复步骤 4～步骤 7，得出多次（如 5 000 次）计算结果。

(9) 将从步骤 8 中获得的 5 000 个值绘制成抽样分布图。

(10) 将经过检验的法则的收益率放入抽样分布，然后计算概率值（法则的收益率的随机分布等于或者大于经过检验的法则的收益率）。

将计算机密集型方法应用于单一法则的回溯测试

这部分内容描述的是两种计算机密集型假设检验方法对单一法则的应用：用道琼斯运输指数（输入序列 4）作为 91 天通道突破法则¹⁴的输入序列，我们把这个法则称为 TT-4-91。而本书第二部分中需要检验的法则，我们将在第 8 章详细介绍，现在的主要任务就是叙述假设检验。法则 TT-4-91 在 1980 ~ 2005 年曾被用于对标准普尔 500 指数发出多头和空头信号。在这段时间里，该法则利用消除趋势的标准普尔 500 指数来计算法则的每日收益率，并获得了 4.84% 的年收益率。在消除趋势的数据中，没有预测力的法则的预期收益率是 0。

自举法与蒙特卡罗置换法都用于对法则不具有预测力的原假设进行检验，但问题是：法则 TT-4-91 所获得的 4.84% 的收益率是否足以对原假设进行反驳呢？

1. 运用自举法检验法则的绩效：怀特真实检验

为了得到平均收益率的抽样分布，以消除趋势的标准普尔 500 指数数据作为初始步骤的具体步骤如下。

（1）将法则的每日收益率进行叠合：由于法则在消除趋势的市场数据中产生了正收益率（+4.84%），所以就要把法则的平均每日收益率（每天的平均收益率大约是 0.019 2）从法则每一天获得的收益率当中减去。通过这种变换我们就会得到一组平均值为 0 的每日收益率，由此便可判断该数据是符合原假设的。注意：不要与消除趋势的标准普尔 500 指数数据相混淆。

（2）对每日收益率进行再抽样：由前一个步骤计算得出的以零为中心的每日收益率是可以被替换的。为了证明自举法的有效性，我们需要进行 6 800 次（在原始样本中观察资料的数量）再抽样。

（3）计算平均收益率：计算 6 800 次经过再抽样的收益率的每日平均收益率。

（4）重复步骤 2 和步骤 3 共计 5 000 次，得出 5 000 个再抽样平均值。

（5）绘制经过自举法得出的再抽样平均值的分布。

（6）将 4.84% 的收益率与抽样分布进行比较，并确定落入等于或者

大于 4.84% 这一年收益率的抽样分布面积。这就需要计算等于或者大于 4.84% 这个收益率的 5 000 个经过自举法获得的平均值。

2. 结果：原假设被反驳——法则确定具有预测力

图 5-14 显示了通过自举法获得的抽样分布，这其中包括经过叠加的法则 TT-4-91 的实际绩效。概率值 0.069 2 表明，在 5 000 个经过自举法获得的平均值当中等于或者大于 4.84% 的概率是 0.069。这也就意味着，如果法则的预期收益率确实等于或者大于 0，那么该法则因为偶然（抽样变异性）因素获得等于或者大于 4.84% 的收益率的概率只有 0.07。统计学家把这种结果称为可能显著。我们将在第 6 章中看到，为了发现某种法则下的研究总和，不论该法则的发现是否源于大量的回溯测试检验法则，都可以影响到概率值和法则的显著性。我们此处引用的结果假设这种法则是用于回溯测试的唯一法则。

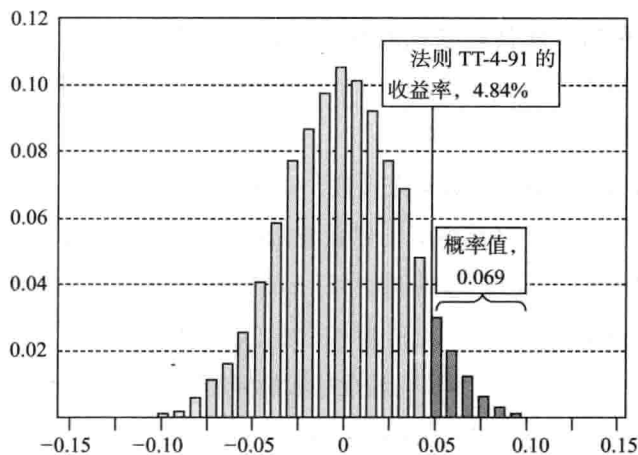


图 5-14 经过自举法得到的法则 TT-4-91 的抽样分布和概率值

3. 运用蒙特卡罗置换法检验法则的绩效

应用蒙特卡罗置换法对法则的回溯测试绩效进行检验，其具体步骤如下。

(1) 法则的每日输出值的时间序列将在其进行回溯测试阶段产生的正常序列中展开，即有 6 800 个这样的值 (+1 和 -1 的排列顺序)。

(2) 对这 6 800 个数值前一天的价格变化进行消除趋势后得到的值写

在一个小球上,并将这 6 800 个小球放入容器中。

(3) 搅动容器,然后从中随机抽取一个小球,并与独立的法则输出值进行配对。因此,每一个法则的每日输出值都会与一个独立的前一天的标准普尔 500 指数的收益率进行配对。在此,我们不需要使用替代法,直到所有的小球与法则的输出值完成配对。由于市场的每日收益率与法则的输出值都是 6 800 个,所以当配对完成以后,容器就会变成空的。

(4) 把每一个法则的值(+1 或者 -1)乘以标准普尔 500 指数的收益率,就会得出下一个交易日通过噪声规则获得的收益率。通过这个步骤,可以产生 6 800 个法则的每日收益率。

(5) 计算从步骤 4 中得到的 6 800 个值的平均值。这个平均值就是噪声规则得出的第一个蒙特卡罗置换法的平均收益率。

(6) 重复步骤 3 ~ 步骤 5,共计 5 000 次。

(7) 绘制 5 000 个蒙特卡罗平均值的抽样分布图。

(8) 将法则 TT-4-91 的收益率与抽样分布的收益率进行比较,然后确定蒙特卡罗平均值等于或者大于法则 TT-4-91 的收益率的概率值。

4. 结果:该法则可能具有预测力

事实证明,由自举法与蒙特卡罗置换法得出的概率值基本相同。

估算

估算是统计推断的另外一种形式。与检验假设不同的是,估算的目的是估计总体参数的大概值,而假设检验的目的则是接受或者排斥总体参数的数值。在我们的案例中,估算被用来对某一法则的预期收益率进行估计。

点估计

估计方法有两种:点估计和区间估计。点估计是指接近总体参数的单一值,如某法则的预期收益率是 10%;区间估计是指总体参数位于给定概率水平的一组数值。我们来看下面的表述:某法则在 0.95 的概率水平上的预期收益率范围在 5% ~ 15%。

实际上,我们已经做过点估计,只是没有像现在这样描述过。每一次我们计算抽样平均值和估计总体平均值的时候,都要用到点估计,然而这一事实却为我们所忽视。我们通常用到的点估计量有:平均值、中值、标准差及方差。估计值是通过从总体中提取的观察资料抽样计算得出的。换句话说,点估计就是抽样统计量。计算抽样平均值的公式如图 5-15 所示。

变量 X 的样本平均值

$$\bar{X} = \frac{X_1 + X_2 + X_3 + \dots + X_n}{n}$$
$$\bar{X} = \frac{\sum_{i=1}^n X_i}{n}$$

式中, X_i 表示在变量 X 中独立的观察资料

图 5-15 变量 X 的样本平均值

均值(平均值)在技术分析领域中的应用无处不在,这已经成为了不争的事实,然而,事实证明样本平均值是一种既简洁又有力的估计量。从某种意义上讲,单一值可以对总体平均值做出最好(信息丰富)的估计。这个事实非常重要。

当我们在分析用于判断估计量资质的标准时,如何才能使样本平均值包含的丰富信息更加清楚呢?估计量的判断标准有:无偏性、一致性、有效性及充足性。运用这 4 种标准,可以证明样本平均值是总体平均值的最佳估计量。¹⁵

如果预期收益值等于总体的值,我们就称其为无偏差性。换言之,如果估计量与总体实际的值之间没有偏差,那么它的平均值为 0。也就是说,样本的平均值与总体平均值之间是不存在偏差的。因此可以说,历史样本的平均收益率与即将实现的预期中的收益率之间是没有偏差的。

估计量的另外一个标准是一致性。随着样本容量的扩大,如果估计量的值与总体参数的值趋于一致,那么我们就说这个估计量是一致的。大数法则告诉我们,样本平均值就是这样的。

估计量的有效性,这一标准与抽样分布的宽度有关。正如前面提到的那样,估计量就是抽样统计量,所以它会有抽样分布。能够产生最窄抽样分布的估计量就是最有效的估计量。换句话说,最有效的估计量具有最小的标准误差。¹⁶ 样本平均值与样本中值对于对称分布的总体来说是没有偏差的,对总体平均值也是一致估计的。然而,样本平均值要比样本中值更加有效。对大样本来说,平均值的标准误差要比样本中值的标准误差少

80%。¹⁷

估计量的最后一个标准就是它的充足性。如果能够利用所有可以获得的样本数据，而其他的估计量又无法添加任何有关参数的信息，我们就说这个估计量是充足的。¹⁸从某种程度上讲，样本的平均值具有充足性。

区间估计——置信区间

比点估计包含更多信息的是区间估计，也称作置信区间。我们将在下面的内容中介绍。

1. 置信区间告诉了我们什么

点估计具有局限性，因为它表达了由于抽样误差所导致的不确定性，而置信区间则通过将点估计与有关估计量抽样分布的信息相结合的方法来解决这个问题。

置信区间是指围绕在点估计周围的一系列数值。置信区间是由上限值和下限值的两条边界来定义的。此外，置信区间通过概率告诉我们真实的总体参数值落入置信区间边界的可能性。因此，对平均值来说，在 90% 的置信水平上，位于边界范围之内真实的总体平均值的概率是 0.90。

当考虑置信区间能够告诉我们什么的时候，最好想一想如果要构成一个像 90% 这样数量庞大的置信区间，而且每一个都是基于从总体中提取出来的独立的数据样本时，会发生什么情况？如果我们真的这么做了，那么 90% 的区间就会真的围绕在总体参数值的周围。展开来讲，大约 10% 的置信区间没有包括在总体参数之内。10% 的置信区间如图 5-16 所示。因为我要保持图形的简单明了，所以采用了比较小的数值。真实的、但是未知的总体平均值用希腊字母 μ 表示。样本平均值由在每个置信区间中的圆点表示。请注意有一个置信区间没有包含在总体平均值内。

研究者可以选择自己需要的置信区间，如 99% 的置信水平等于概率为 0.99 的真实总体平均值。当然，对更高水平的置信区间要设定一定的数值，即这个区间要更宽。换句话说，置信水平设定的越高，就越会减少精确度。在图 5-17 中，将置信水平增加到 99% 的水平，那么在图 5-16 中由 90% 的置信区间造成的错误就会被消除。减少的错误率是通过运用更宽（不精确

的)的置信区间来完成的。

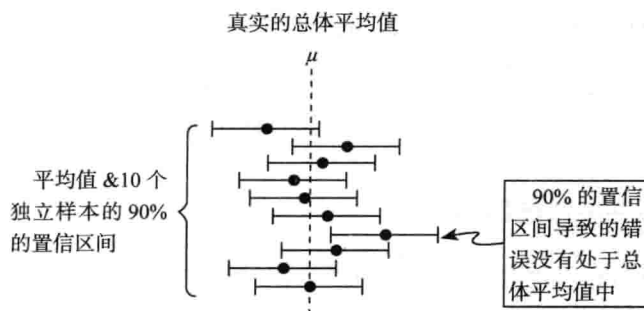


图 5-16 90% 的置信区间 (0.90 的概率是正确的)

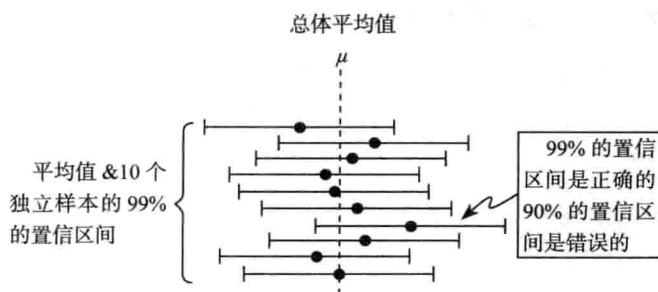


图 5-17 99% 的置信区间 (0.99 的概率是正确的)

如果 100 个独立的数据样本要对某法则进行回溯测试, 并且 90% 的置信区间都被绘制在每一个样本中被观察的平均收益率周围, 那么这些置信区间中的大约 90% 就等于法则在总体中的真实收益率或者预期收益率。图 5-18 显示, 在回溯测试中获得 7% 的收益率的法则, 它的置信区间是 90%。我们可以确定, 在 0.90 的概率水平上, 该法则的预期收益率在 2% ~ 12%。

2. 置信区间与抽样分布的联系

置信区间来源于与计算假设检验概率值相同的抽样分布。我们已经学习了有关抽样误差 (抽样变异性) 的知识, 可以说样本的平均值等于未知的总体平均值, 加上或者减去它的抽样误差。它们之间的关系请参见图 5-19 中最上面的公式。通过对上面这个公式中的各项进行重新排序, 就得到了图 5-19 中最下面的公式。也就是说, 总体参数的平均值等于已知的样本的平均值加上或者减去抽样错误。

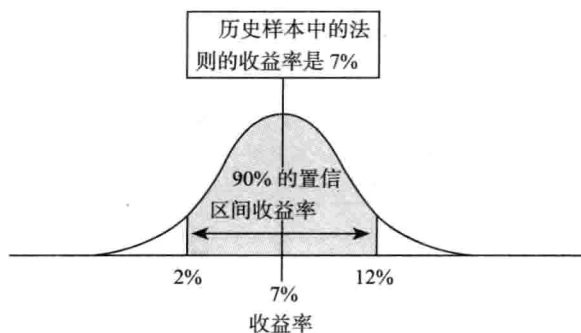


图 5-18 进行回溯测试的法则的 90% 的置信区间

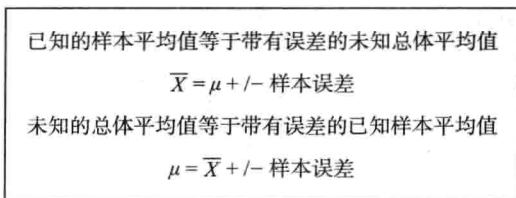


图 5-19 已知的样本平均值等于带有误差的未知总体平均值

最下面的公式说明, 尽管我们不知道精确的总体平均值, 我们可以利用已知的样本平均值以及平均值的抽样分布来获得一组数值, 这些值等于在具体概率水平上的总体平均值。在实操方面, 如果我们要重复下面的步骤 1 000 次——计算样本平均值和 90% 的置信区间, 总体平均值就会落入 0.9 的置信区间, 如图 5-20 所示。图 5-20 中的步骤只是重复了 10 次。注意图 5-20 中 10 个置信区间中有一个没有包括在总体平均值当中。这一部分的知识点主要是, 置信区间的宽度来源于根据样本抽样的宽度。

正如我们之前说过的那样, 置信区间是以同样用于假设检验的抽样分布为基础的。然而, 就置信区间而言, 抽样分布只是简单地把它在假设检验当中所占有的位置进行了位移。在检验假设当中, 抽样分布集中在被假设的总体平均值上, 例如, 0; 对置信区间来说, 抽样分布集中在超过样本平均值的位置上, 例如, 7% (见图 5-21)。

3. 用自举法产生置信区间

自举法可以用来推导出置信区间。它的推导程序与生成假设检验的抽样分布的程序是相同的。

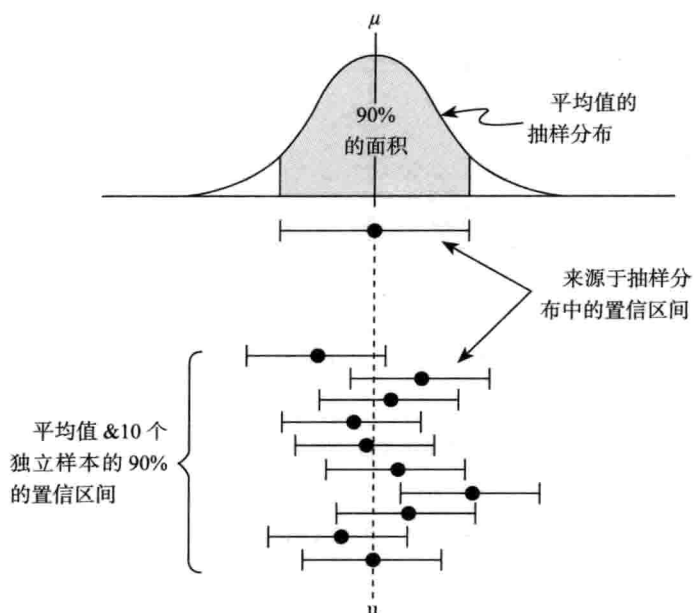


图 5-20 置信区间与抽样分布之间的联系

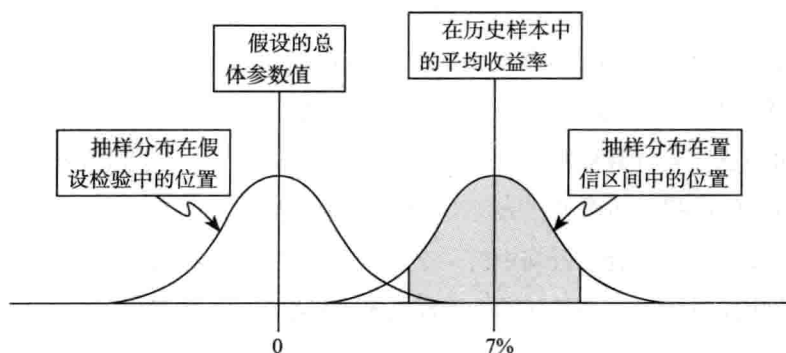


图 5-21 抽样分布在假设检验和置信区间中的位置

目前有很多种方法用于计算自举法的置信区间，我们在这里介绍其中的一种，即“引导百分位数法”。这种方法不仅流行，而且简便易用，经常可以产生很好的结果。而更为复杂的方法则超出了本书的范围。

需要指出的是，蒙特卡罗置换法并不能够产生置信区间。这是因为该方法与估计总体参数的值和检验关于此数值的观点无关。我们之前曾经提过，蒙特卡罗置换法的检验对象是与含有法则信号的信息相关的观点。具

体来说,原假设认为,根据蒙特卡罗置换法,由法则发出的多头头寸和空头头寸是缺少有用的未来市场变化信息的。因为它并没有提及总体参数(如法则的平均收益率),所以也就不可能产生置信区间!

用引导百分位数法创建置信区间的步骤如下:假设某法则的收益率经过 5 000 次的再抽样,并且计算出了每一个再抽样的平均值,那么就会产生 5 000 个不同的再抽样平均收益率。由于抽样变异性的缘故,这些平均值也会不同。接下来我们再假设将这 5 000 个平均值按照从高到低的顺序进行排列。然后,根据对置信区间的需求,将最高的百分比 x 和最低的百分比 x 从排序当中移除,就得到:

$$x = \frac{100 - \text{置信区间}}{2}$$

如果我们需要 90% 的置信区间,就要从 5 000 个再抽样平均值当中将最高的 5% 和最低的 5% 移除,也就是说要移除最高的 250 个再抽样平均值和最低的 250 个再抽样平均值。在移除这些极端的数值以后,剩下的再抽样平均值中最高的值就构成了 90% 水平置信区间的上限,而与其相对应的最低的再抽样平均值就构成了 90% 水平置信区间的下限。而在 99% 水平上的置信区间只需要从包含抽样分布的 5 000 个再抽样平均值中移除最高的 25 个(最高值的 5%) 和最低的 25 个(最低值的 5%) 再抽样平均值。

4. 假设检验 VS. 置信区间: 潜在的矛盾

很多精明的读者也许会想到一个问题,即假设检验和置信区间很有可能会得出不同的法则预期收益率。这种想法来源于假设检验集中在抽样分布右边的尾部,而置信区间则集中于抽样分布左边的尾部。这意味着 90% 水平的置信区间的下限有可能表明法则的预期收益率有 5% 的概率小于 0,而假设检验在 0.5 的显著水平上就可以排斥原假设。换句话说,置信区间可以告诉我们该法则不具有预测力,而假设检验则认为该法则没有预测力。从理论上讲,假设检验和置信区间所得出的结论应该是相同的。也就是说,如果 90% 水平的置信区间的下限告诉我们法则的预期收益率小于 0,那么在显著性为 0.5 水平上的假设检验就不会排斥原假设。

当抽样分布不对称时(向右或者向左偏),矛盾的结论就会出现。从

图 5-22 中我们可以清楚地看到抽样分布向左侧偏斜，并且显示在两个位置上。在图 5-22 中下面的那幅图中，抽样分布的位置就好像它要对原假设进行检验一样。因为抽样分布右边的尾部位于进行回溯测试法则的平均收益率之上的部分小于 5%，因此，检验表明该法则的显著性处于 0.5 的水平。换句话说，对原假设的排斥可以被备择假设所替代，因为备择假设认为法则的预期收益率大于 0。

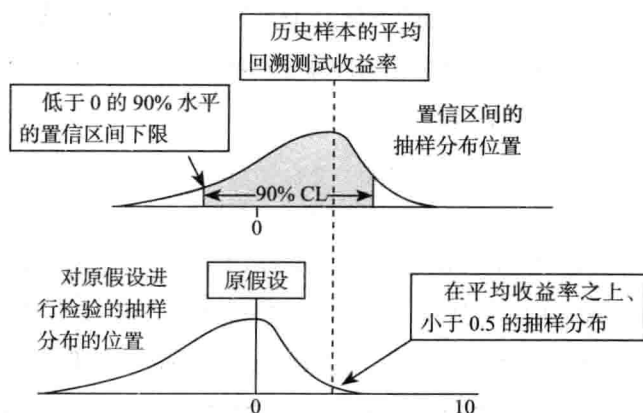


图 5-22 潜在的矛盾结论：假设检验 VS. 置信区间

图 5-22 中上面的那幅图中显示了抽样分布的定位是通过运用引导百分位数法来构成 90% 水平的置信区间。注意置信区间的上限是低于 0 收益率的。这一点说明法则的真实收益率小于 0 的概率超过 0.5。换句话说，90% 水平的置信区间导致了与假设检验得出的结论相反的结果。图 5-22 显示这种歧义是由于抽样分布的不均匀所导致的。

幸运的是，这一问题并没有影响到本书的研究工作，本书采用样本平均值作为绩效统计量。至关重要的中央极限定理可以保证在样本容量足够大的情况下，平均值的抽样分布不会出现严重的偏斜（不均匀）。中央极限定理认为，随着样本容量的增加，平均值的抽样分布会趋于形成不均匀的钟形态。而其他的绩效统计量则不会出现这种情况。在抽样分布不均匀的情况下，可以通过旋转抽样分布的引导技术来缓解这个问题。令人遗憾的是，其他方法都有自身的局限性。不管怎样，这会使得平均收益率得出引人注目的绩效统计量，进而用于对法则的检验。

5. 法则 TT-4-91 的置信区间

这部分内容描述的是法则 TT-4-91 的置信区间的样本。图 5-23 显示了 80% 的置信区间，即叠加在位于通过回溯测试观察到的 +4.84% 收益率处的自举抽样分布。80% 的置信区间的下限是 +0.62，上限是 9.06%。这说明，如果法则 TT-4-91 对 100 个独立的数据样本进行回溯测试，且 80% 的置信区间都位于每一个样本的平均收益率周围，也就是大约 80% 的抽样，那么该法则的真实预期收益率就会被置信区间所包围。

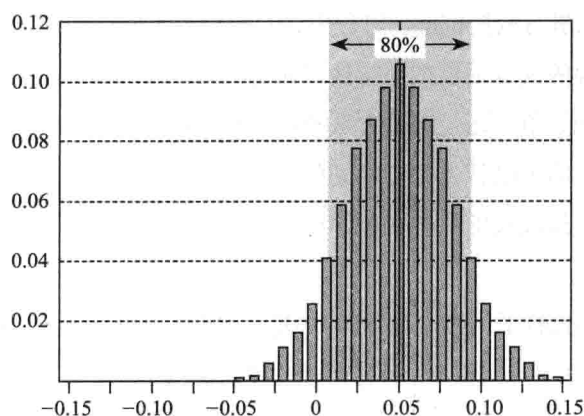


图 5-23 抽样分布以及法则 TT-4-91 的 80% 的置信区间

第6章

Chapter 6

数据挖掘偏差：客观技术分析的黄铜矿

在对法则进行数据挖掘的过程中，我们对众多的法则进行了回溯测试，并且把观察到的具有最佳绩效的法则挑选了出来。也就是说，数据挖掘中存在着绩效的竞争，通过这种方法我们可以把绩效最佳的法则从中挑选出来。但是现在的困难就是把观察到的具有最佳绩效的法则从那些有意夸大其在未来有上佳表现的法则中挑出来。我们把这种系统性误差称为数据挖掘偏差。

尽管存在这样的困难，但是数据挖掘仍然是一种非常有用的研究方法。它可以用数学的方法证明在所有经过检验的法则当中，具有最高绩效的法则很有可能在未来表现出上佳的业绩水平，并且提供充足的观察资料用于计算其性能统计。¹ 换句话说，即使观察到的法则的最佳业绩的确带有明显的偏差，数据挖掘的这种做法也是值得的。本章的内容将会对这种偏差出现的原因、在对最优法则未来的绩效进行推理时把它考虑进去的原因，以及如何进行推理这几方面做出解释。

我首先用一些奇闻轶事对这个抽象的主题进行介绍，它们当中只有一个会与对法则进行的数据挖掘有联系。这些有助于我们对后面的资料的理解。那些一上来就想直奔主题的读者可能会选择跳过“数据挖掘”的这部分内容。

下面的定义将会贯穿本章的内容，为方便读者阅读，我将它们列在此处。

- **预期绩效 (expected performance)**：是指在马上实现的预期当中，某法则的预期收益率。我们也可以将它称为法则的真实绩效，也就是

具有真正预测力的法则的绩效。

- **观察到的绩效 (observed performance)**: 在回溯测试中法则所产生的收益率。
- **数据挖掘偏差 (data-mining bias)**: 指的是所观察到的最佳法则的绩效与其预期绩效之间的预期差异。**预期差异**是指它是从衡量已经观察到的最优法则的收益率与该法则实际的预期收益率之间差异的大量实验当中获得的。
- **数据挖掘 (data-mining)**: 是在规模较大的统计数据库中寻找形态、模型、具有预测力的法则的过程。
- **最佳法则 (best rule)**: 在对法则进行回溯测试的过程中得到的绩效进行比较之后, 获得最佳绩效的那种法则。
- **样本内数据 (in-sample data)**: 用于数据挖掘的数据 (例如, 法则的回溯测试)。
- **样本外数据 (out-of-sample data)**: 不用于数据挖掘或者回溯测试过程中的数据。
- **法则集合 (rule universe)**: 在数据挖掘投资中所有进行回溯测试的法则的集合。
- **全域规模 (universe size)**: 包含在法则集合当中的法则的数量。

✓ 落入陷阱：数据挖掘偏差的故事

下面的故事纯属虚构。在我从事统计学研究的若干年前, 有一位主办商找到我, 说他正在为一场商业投资展示会寻找赞助, 并且告诉我这场展示将会得到丰厚的利润。这场展示会的主要内容就是一只猴子可以通过在文字处理器的键盘上跳舞的方式写出莎士比亚的诗歌。

在每一次的展示中, 这只被主办方称为“诗人”的猴子就会被放在键盘上, 并且通过大屏幕向观众显示它的杰作。可以肯定的一点是人们蜂拥而至, 都想一睹“诗人”的风采, 我通过卖票获得的收益不仅能把投资的 50 000 美元挣回来, 而且还能大赚一笔。至少主办方是这么说的。他们不

断重复：“千万不要错过！”

我被这件事吸引住了，但是想证明这只猴子是否真的能够写出莎士比亚的诗歌。任何投资者都会这样做。我勉强算是对此进行证明了吧。在一封算不上安慰信的信件中，他们说，“经过对‘诗人’以前作品的分析，可以得出它确实能够写出‘活着还是死去，这是一个问题’（to be or not to be, that is a question）这些词。然而，我们并不知道这些词是在什么样的环境下写出来的”。

我现在最需要的就是在现场证明这一切。令人遗憾的是，我的要求并没有得到满足。主办方不耐烦地解释说，猴子喜怒无常，而且现在还有许多其他的投资者吵着要购买有限的票。所以，我想抓住这个机会并投入 50 000 美元。我非常自信地认为，钞票源源不断地流入我的钱包只是时间的问题了。

第一次展示的时刻到来了。卡耐基音乐厅内人头攒动，大家迫不及待地等待着第一个单词的出现。人们把目光全部集中于大屏幕之上，“诗人”的第一行作品出现了。即：

Lkasldlk5jf wo44iuldjs sk0ek 123pwkdzsdidip'adipjasdopiksd

接下来的事情变得越来越糟糕了。观众们大声叫嚷着要求退票，而“诗人”也开始在键盘上手舞足蹈起来，然后蹦蹦跳跳地下了台。此时，我不得不放弃了自己的念头，我知道自己的钱彻底打了水漂儿。

到底发生了什么？这时候想这些事情已经不重要了。主办方并没有透露一件重要的事实。“诗人”是从 1 125 385 只猴子中挑选出来的，在过去的 11 年 4 个月零 5 天的时间里，这些猴子每天都能获得在键盘上跳舞的机会。一台计算机将这些莫名其妙的字母组合成单词，并且把这些单词与莎士比亚著作中的语句进行匹配。而这个“诗人”是第一只完成这项任务的猴子。

即使是在我缺少统计学知识的时候，我就已经对我的投资产生了质疑。仅仅从常识来判断，我就知道在主办方大量的吹嘘当中，他们所引用的莎士比亚诗歌都是偶然出现的。对一群由猴子创造的数万亿个字母进行自由的数据挖掘，提高了幸运的字母序列成为一个虚拟的确定性的概率。“诗

人”根本就不会写作，只是一切纯属巧合而已。

亲爱的布鲁特斯，他的错误在于他是一个非常坦率诚实，但是在统计学方面又是如此的一无所知的主办者。他完全被数据挖掘偏差蒙蔽了，而这些显著性又可以归结为由数据挖掘所导致的结果。尽管我遭受了损失，但是我并不想苛刻地去评价主办商。他坚定地相信他发现了一只神奇的猴子。他被一种缺少对统计学和概率进行评估的直觉误导了。

顺便说一句，主办商一直把这位“诗人”当作宠物，并且让它继续在那个曾经神奇无比的键盘上跳舞，希望可以得到新的文学上的实证。但是为了养家糊口，他现在正在出售根据类似的方法开发的技术培训系统。他用大量的会跳舞的猴子来开发规则，有些规则在历史数据中的表现非常好。

用棒球统计来证明上帝的存在

运动统计学的采集者也受到了数据挖掘偏差的迷惑。例如，一个名叫诺曼·布鲁姆（Norman Bloom）的人得出了这样一个结论，即根据从棒球统计法中发现的令人感兴趣和不寻常的形态，可以证明上帝的存在。在对他的数据库进行数千次的研究之后，这位专注的数据挖掘器发现，他所坚信不疑的形态是如此的令人惊奇，只能用由上帝创造的宇宙秩序来解释。

其中一种由布鲁姆所述的形态如下：乔治·布雷特（George Brett）是堪萨斯城队的三垒，他在季后赛第三场比赛中凭借着本人的第三记本垒打帮助球队把总比分扳为3-3。根据布鲁姆的推理，数字3和许多事情都有联系，这肯定是上帝一手安排的。另外一个由他发现的有趣的形态是关于股票市场的，即道琼斯工业平均指数在1976年一共有13次突破了1000点大关，这与美国在1776年由最初的13个殖民地组成联邦这一事实惊人的相似。

罗纳德·卡恩（Ronald Kahn）²指出，布鲁姆在他做出未经证明的结论这一过程中犯了很多的错误。首先，他并没有理解随机的作用，实际上，那些看上去比较罕见的巧合如果经过充分的探索是完全有可能出现的。布鲁姆通过对上千种可能出现的联系的评估得出了他所说的神奇形态。布鲁

姆在他开始研究之前，并没有明确描绘出任何重要的形态。相反地，他使用的是在事实之后再做定义的这一约定俗成的标准。任何对他这种有趣的和不寻常的构想进行的攻击行为，都会被认为是非常重要的。卡恩指出，当人们用一种杂乱无章的方式来进行研究的时候，他们肯定会发现“有趣的”形态。

发现隐藏在“旧约”当中的预测

即使是研究《圣经》的学者也会落入数据挖掘的陷阱。在下面的例子中，那些用心良苦而又对统计学一无所知的研究者发现，对世界上发生的主要事件的预测是可以通过将旧约当中的文字翻译成编码的形式来做到的。未来的知识暗示了对文字的灵感来源于无所不知的创造者。然而，对于圣经密码存在着一些问题。它们总是在被预测的事件发生之后才能够发现这一切。换句话说，圣经密码具有后见之明的特点。³

圣经密码是一种集群，它是把按照特定的距离或者按照特定的数量隔开的字母组合在一起的文字。我们把这种具有创造性的文字称为等距离字母序列（equal letter sequences, ELS）。密码的研究者可以随意确定间隔区间的距离，并把包含在集群内的单词按任意的形态进行排列，但前提是这个集群要出现在被研究者压缩了的原文的范围之内。⁴至于这个被压缩的范围是多少，以及会做出什么样的预测，都要等到这个密码被发现之后才能确定。请注意由事实之后定义的评估并不符合科学的标准。

圣经密码的研究者争辩说，密码的出现是不具有统计学意义的，它只能通过由上帝安排在这里的密码来解释。这些天真的数据挖掘器所犯的基础错误在于他们没有理解通过足够多的研究（数据挖掘）得到这种形态的概率也是很高的这一事实。因此，研究者发现这些密码可以和人名、地点以及历史上的重要事件相对应。例如，1990这个词和处于同样范围内的“萨达姆·侯赛因”“战争”这些词在一起，就不是需要用形而上学来解释的小概率事件了。然而，在1992年第一次伊拉克战争爆发之后，单词的组合形态似乎预测到了1990年的伊拉克战争。请记住下面的单词：“布什”“伊拉克”“入侵”以及“沙漠风暴”，很明显，这些单词都出现在了

1990 年的伊拉克战争的预测当中。实际上，有很多的单词组合可以反映出发生在 1990 年的战争，当然，这些都是在特定的历史事件发生之后才发现的。

迈克尔·卓斯宁 (Michael Drosnin) 是《华尔街日报》和《华盛顿邮报》的记者，并没有受过统计学方面的专业训练。在他 1997 年出版的《圣经密码》(Bible Code) 一书中，他对埃利亚胡·瑞普斯 (Eliyahu Rips) 博士的研究做出了描述。瑞普斯博士是群理论数学方面的专家，这一理论是与数据挖掘偏差并无太大关系的数学分支。尽管卓斯宁声称圣经密码得到了一些著名的数学家的认可，但是有 45 位统计学家在对瑞普斯的作品进行检验之后，认为他的说法并没有足够的说服力。⁵ 数据挖掘偏差的本质问题就是错误的统计学推断。

统计学家们不赞成圣经密码研究者们进行的毫无限制的探索。因为他们犯了一个被称作“过度燃烧的自由度”的数学错误。对于精通统计学的人来说，犯这种错误是违反规则的。就像巴里·西蒙 (Barry Simon) 博士指出的那样，“用怀疑的眼光看待圣经密码”，⁶ 在 *Chumash* (摩西五经中的一本) 中，其中的 14 卷包含了圣经的内容，假设我们把等距离字母序列 (ELS) 的间隔区间设定为 1 ~ 5 000 时，就可以产生大约 30 亿个由这部分原文所组成的可能出现的单词。这种根据兴趣将那些自创的单词随意组合的行为，与在键盘上跳舞的猴子所写出来的数以吨计的莫名其妙的字当中做研究的行为是一样的。

研究圣经的学者们寻找计算法则的这一愚蠢行径是显而易见的，因为当他们把这一套方法应用在诸如《西尔斯分类目录》《白鲸》、托尔斯泰的《战争与和平》以及芝加哥市的电话号码簿等非圣经的文本中时，这种情况就发生了。当对这些文本进行搜索的时候，事后对历史事件进行预测的密码也可以被找到。这说明了密码只是搜索方法的副产品，而不是要搜索的文本产生的密码。

在卓斯宁最近出版的《圣经密码 II》一书中，他在一篇名为“倒计时”的文章里描述了密码是如何预测 2001 年“9·11”恐怖袭击事件的。你也许会问他为什么没有在恐怖袭击事件发生之前向我们发出警告呢？他没有

这样做是因为他根本就不可能这么做。他是在悲剧发生之后才发现这些密码的。卓斯宁成为了那些用心良苦，而又被数据挖掘偏差所欺骗的天真的研究者中的又一个实例。

出现在新闻记者当中的数据挖掘

新闻记者也会受到数据挖掘偏差的蒙蔽。在 20 世纪 80 年代中期，他们报道了一则关于伊夫林·亚当斯 (Evelyn Adams) 的故事，亚当斯在 4 个月的时间里赢得了两次新泽西州彩票大奖。⁷ 报纸认为这种事情发生的概率是 17 万亿分之一。实际上，找到赢得两次大奖的人很容易，这个故事不像记者们想的那样具有更高的新闻价值。

亚当斯先生或者其他任何人两次赢得彩票大奖的事前概率是 17 万亿分之一。然而，通过寻找所有已经赢得过两次大奖的人的事后概率却是非常高的。哈佛大学的统计学家珀西·迪亚科尼斯 (Percy Diaconis) 和弗雷德里克·莫斯特勒 (Frederick Mosteller) 对此估计的比率是 1/30。

符合事后概率是最关键的环节。它是指在结果已经知晓之后对数据进行仔细的分析。就像任意一只猴子在未来可以创作出莎士比亚的著作的概率是很小的一样，这种概率只会发生在数百万只猴子中的某些猴子身上，它们能够创作出诗歌的概率要高得多。只要机会足够的多，那么通过随机就可以得到意外的结果。这种事情看上去几乎不会发生，但实际上它却极有可能发生。

挖掘联合国的数据库寻找黄金和黄油

大卫 J. 莱因韦伯 (David J. Leinweber) 是加州理工学院的教授和定量养老金管理公司 First Quadrant 的全权合伙人，他曾经对金融市场的研究者发出过有关数据挖掘偏差的警告。为了揭示过度搜索的误区，他测试了联合国数据库当中的几百个经济时间序列，以此来发现一个对标准普尔 500 指数具有高度预测性的相关序列。结果是，他找到的与标准普尔 500 指数相关度最高的就是孟加拉国的黄油生产水平，即 0.70，这对于经济预测领域而言是相当高的了。

仅从直觉上来说，孟加拉国的黄油和标准普尔 500 指数之间的高度相关貌似是真实的，但是现在试想一下，如果最高相关度的时间序列与标准普尔 500 指数有这样一种貌似合理的联系，直觉就不会警告我们了。正如莱因韦伯所指出的那样，当把所有的时间序列都考虑在内的时候，孟加拉国的黄油和准普尔 500 指数之间的联系就没有统计的显著性了。

从本质上来说，不论是运动统计学的探索者、圣经、随机写作的猴子、买彩票的人，还是金融市场历史，如果不把数据挖掘偏差考虑进去的话，那么就会导致错误的结论。

在客观的技术分析中出现的错误知识的问题

技术分析包括两个相互排斥的领域，即主观的技术分析和客观的技术分析。现在回想一下，客观的技术分析方法仅限于那些定义清楚的，可以被简化为计算机算法，并且可以进行回溯测试的方法。而另外一种方法就是主观的技术分析。

这两个领域都会受到错误信息的侵害，只是谬误的表现形式各不相同。主观的技术分析松散的定义并没有包含太多的信息，所以就会产生不能被检验的预测，⁸ 因此，它就不会受到任何实证的质疑。如果没有删除那些没有价值的思想这一重要过程，谬误就会不断地积累。结果是，主观的技术分析不是一个正规的知识体系，而是基于奇闻逸事和直觉这一薄弱基础之上的民间传说的集合。

只要在随机性（抽样变异性）和数据挖掘偏差当中把回溯测试的结果考虑进去，客观的技术分析方法就会具有成为有效知识的潜力。由于许多客观技术分析的实践者并不关注这些效果，所以谬误也会在这一领域内不断积累。

错误的客观技术分析表明样本外的⁹绩效会出现下降，即一种法则在用于回溯测试的抽样当中的绩效表现尚可，而在样本外的数据中的绩效则表现糟糕。这个问题如图 6-1 所示。

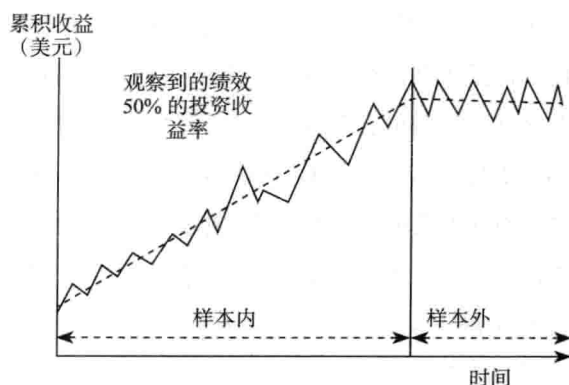


图 6-1 样本外的绩效出现下降

对样本外的绩效出现下降的解释：旧和新

样本外的绩效出现下降是个众人皆知的问题。¹⁰ 客观的分析师为此做出了一些解释。而我则在数据挖掘偏差的基础上提出了一个相对比较新的解释。

对这个问题貌似合理的解释要归结为随机变量。尽管由于抽样变异性的原因，使得某种法则的绩效在不同的历史抽样当中的表现不尽相同，但是这种解释并不能就此等同于实证。如果随机变量是可靠的，那么样本外的绩效就会高于样本内绩效，尽管通常情况下它是比较低的。任何经历过回溯测试的人都知道样本外的绩效在更多的时间里是低于样本内绩效的。

第二个理论依据是关于当处于样本外的时间时，市场的动态变动。我们可以合理地假设金融市场是一个非恒定系统。¹¹ 然而，对法则每一次出现在样本外的失败都是由于市场动态变动造成的这种假设是不合理的。对于当法则的移动趋势从技术人员的实验室里转换到现实的交易当中的时候总是发生变化这种情况而言，这实在是让人觉得又奇怪，而又不可能实现的事情。我们对市场的动态变动与法则在样本外失败的频率是一样多的假设是不合理的。

以不断变化的市场动态为基础，我们又可以做出另外一种非常复杂的解释，在此我们引用额外的假设，即变化的发生是因为法则已经被很多的交易者所采用。也就是说，用该法则进行买卖的交易者已经将市场的形态

破坏了，而这种形态是要对法则的回溯测试产生的盈利性进行说明的。这种推理还缺少合理性。假设现在有无数的法则需要进行解释，而让足够多的参与者采用任何一种法则又是不可能的。即使当很多的交易者采用相类似的法则，即与期货交易基金一样，采取客观的趋势跟踪方法，那么由此降低的法则的绩效似乎是归结于市场波动引起的变化，而与所使用的技术系统¹²没有关系。

我们还可以在随机性¹³的基础上对样本外的绩效的失败做出另外一种解释。这种解释认为，市场的动态是系统的或者模式化的行为与随机性或者噪声的结合体。有效的法则利用市场的动态特征。因为系统化的部分应当在样本外的数据中持续显示，这一点是法则获得盈利的可靠来源。相反，随机成分是无法复制的现象，它在每一个数据抽样当中的表现都不相同。法则在样本内的绩效存在部分运气的成分，即在法则发出的信号和市场的非重复性噪声之间存在着对应的叠合。这些依靠运气成分获得的利润是无法体现在样本外的，因此，未来的绩效将会低于之前的绩效。以上的解释已经非常接近我们的目标了，但它还是不完整的。随机性只是其中的一个负面因素。

对样本外的绩效的下降的更加完整的描述是以数据挖掘偏差为基础的。它有两个负面因素：①随机性，是观察的绩效中相对较大的部分；②数据挖掘的逻辑，是指在所有进行回溯测试的法则产生的绩效当中挑选出具有最佳绩效的法则，以供数据挖掘器的检验之用。当这两种作用相结合的时候，它们会导致观察到的最佳法则的绩效超过其未来的（预期的）绩效。因此，最佳法则未来的绩效很有可能会比导致其从被检验的法则当中挑选出来的绩效水平差得多。

数据挖掘偏差对样本外的绩效的下降的解释要比基于变动的市场动态的解释更有优势。这两种解释都符合样本外的绩效趋于下降的事实。然而，后一种解释引用了市场动态已经发生变化的假设。而数据挖掘偏差则没有引用这一假设。简单来说，数据挖掘的过程证实了在进行回溯测试的过程中，对法则的挑选受到了运气成分的影响。当在一组善于解释某种现象的法则中进行选择的时候，最明智的做法就是选择那个能够做出最少假设的

法则。我们把这种简化的原理称为“奥卡姆剃刀”，这一概念我们曾在第3章中进行过讨论。

在接下来的内容里，经过观察的法则的绩效可以分解为两个独立的部分，即随机性和法则内在的预测力。在这两个部分中，随机性是到目前为止最重要的部分。因此，具有最佳观察绩效的法则最有可能受到好运气的眷顾。也就是说，随机增加的超过标准水平的法则的绩效要归结于它自身真实的预测力，前提是如果它真实的拥有的话。与其相反的是拥有较差绩效的法则。这很有可能是运气产生的负面影响。通过选择观察到的最佳绩效的法则，数据挖掘器最终挑选了一个好运气占大部分的法则。由于不能指望运气会每次都眷顾，因此就要把该法则的样本外绩效回归到能够代表它内在的真实预测力的水平上去。这就有可能使得样本外的绩效低于绩效的水平，从而让该法则赢得绩效间的竞争。从本质上说，最佳法则的样本外绩效的下降很有可能会从一个不切实际的高预期值回落，而不是法则本身的预测力出现下降。

数据挖掘

数据挖掘是以形态、法则、模型、函数的方式从大型的数据库当中对知识进行提取。当某种知识包含多元变量、非线性关系以及高水平的噪声时，人类智慧的局限性和偏差会使得这项工作在无助的思想面前变成了几乎不可能完成的任务。因此，数据挖掘依靠的是计算机算法来进行知识的提取。有关用于数据挖掘的主要方法的讨论，请参见 *Elements of Statistical Learning*,¹⁴ *Predictive Data Mining: A Practical Guide*,¹⁵ 和 *Data Mining: Practical Machine Learning Tools and Techniques*。¹⁶

将数据挖掘比作多重比较过程

数据挖掘是以一种被称为多重比较过程¹⁷（MCP）的方法为基础解决问题的。MCP的基本思想是对不同的结论进行检验，并且挑选出一个符合最佳标准的结论。应用MCP需要3种因素：①定义明确的问题；②一组

备用的结论；③品质因数或者对每一个备用结论的优点进行量化的打分函数。当把所有数据都记录下来之后，将它们进行比较，并且把得分最高的（观察到的最佳绩效）那个备选结论提取出来，作为这个问题的最佳解决方案。

下面我们来分析如何将这种解决问题的模式应用到对法则的数据挖掘中来。

（1）问题：在金融市场中设定多头和空头头寸，并据此获得利润。

（2）一组备用的结论（结论的集合或者解决方案的空间）：由客观的技术分析师提出的一组法则。

（3）品质因数：对诸如历史测试周期、夏普比率以及溃疡指数（Ulcer Index）¹⁸的平均收益率之类的财务业绩进行衡量。

所有法则的绩效是通过回溯测试确定的，我们把获得最高绩效的法则挑选出来。

将对法则的数据挖掘视为规范的搜索

我们可以把数据挖掘看作一种规范的搜索。也就是说，它是寻找产生最佳绩效的法则的标准。这个标准是一组数学/逻辑运算符，它应用于单一的或者是更多的市场数据序列，并且通过法则的指令把它们转变成市场头寸的时间序列。

假设下面经过定义的法则是具有最佳绩效的法则：

如果用道琼斯运输平均指数的收盘价除以标准普尔 500 指数的收盘价的比率大于 50 日移动平均线，那么就应当持有标准普尔 500 指数的多头头寸，反之，则持有空头头寸。

上面的法则包括两个数学运算符号，即比率（ratio）和移动平均（moving average）以及两个逻辑符号——不等式运算符“大于”和表示条件的“如果、那么、其他”，一个数值常数：50，以及两个数据序列：道琼斯运输指数和标准普尔 500 指数。数据挖掘发现，这组标准比其他任何一组经过检验的可替代的标准的绩效都要好。

搜索的类型

数据挖掘的搜索范围由浅入深。它们的不同点之一在于搜索全域的设定宽度。下面的内容就要对3种搜索全域的界定进行分析,我们按照从最狭义到最广义的顺序进行分析。数据挖掘使用的所有方法,不管用多么先进的方法进行简单搜索还是复杂搜索,都会产生数据挖掘偏差。

1. 参数最优化

数据挖掘中最狭义的形式就是**参数最优化**(parameter optimization)。这里所提到的搜索全域只限于除参数值不同以外的相同形式。因此,搜索只局限于发现在特定形式中具有最佳绩效的法则的参数值。

这种形式的法则的实例就是**双重移动平均线交叉法则**(dual-moving-average crossover rule)。它由两个特定的参数组成,即回看期(look-back periods)内的短期和长期移动平均线。当短期移动平均线向上(买入)或向下(卖出)和长期移动平均线交叉的时候,信号就会发出。参数最优化要搜索的就是一组产生最佳绩效的参数值。

能够搜索到的双重移动平均线交叉法则的最大数字等于对短期参数进行检验得出的数值数量与对长期参数进行检验得出的数值数量的乘积。把所有这些组合全部考虑进来的分析指的是详细彻底的,或者是强力的搜索。

目前还有更加智能的搜索方法,这种方法只限于对能够产生好的结果的组合进行搜索。这一目标是通过使用在早期的搜索当中发现能够引导在后期进行检验的参数组合的绩效来实现的。这种更加智能的搜索方法的代表是**遗传算法**(genetic algorithm),这种技术是以生物进化原理为基础的。它的论证能力表现在当可能的参数组合的数量很高时,发现能够以相对较快的速度发挥出理想作用的参数组合。当绩效受到在不切实际的算法基础上提出的更加传统的搜索方法所带来的强烈的随机性的影响时,遗传算法也被证明是有效的方法。对优化算法的出色探讨请参见帕尔杜(Pardo)、¹⁹卡茨(Katz)、麦考密克(Mccormick)²⁰和考夫曼(Kaufman)²¹所写的对各种先进方法的评论。

2. 寻找法则

数据挖掘的广义版本称为寻找法则 (rule searching)。这里提到的法则的全域的不同之处在于它们的概念形式和参数值。双重移动平均线交叉法则就是趋势跟踪的形式体系。因此，在趋势跟踪法则的广义分类中（包括通道突破、移动平均线等），它属于简单的形式。趋势跟踪法则的一般分类就是指包含在其他逆势（平均反转 (mean-reversion)）法则、²² 极值法则 (extreme value rules)、背离法则 (divergence rules)²³ 等分类中的一种技术分析。每一种分类都可以通过特定的法则形式来体现。

本书的第二部分将会在寻找法则的基础上进行数据挖掘的案例研究。研究主要集中在 3 种法则的分类或者主题上：趋势、极端情况和转换以及背离。每一个主题都会通过特定的法则形式表现，这些内容将在第 8 章中进行定义。

尽管寻找法则考虑了大量的法则形式，但是每一种法则的复杂性在搜索的过程中都保留了下来。这里的复杂性是指需要对法则进行说明的参数数量。换成另外一种说法就是，这里对寻找法则的定义，并不包括将简单的法则进行合并之后产生更多复杂的法则。

3. 用可变的复杂性归纳法则

数据挖掘中最为广义和最为突出的形态就是法则归纳法。我们在这里把没有被定义的法则的复杂性也考虑进来。一种复杂的法则有可能被认为是由逻辑运算符叠加的，或者是在多元变量模型中通过与数学函数合并而成的更加简单的法则的合成函数。这种对法则的归纳法进行的数据挖掘的目的，就是发现理想的法则的复杂性。

与在数据挖掘中并不突出的形式相反的是，法则的复杂性是在搜索开始的时候进行定义的，法则的归纳法利用机器学习法（自主归纳法）找出产生最佳绩效的法则的复杂性程度。法则归纳法以检验单一的法则为起点。接下来，对两个法则进行分析，看看它们的绩效是否比单一法则最好的绩效好。再下一步，逐渐地对更多复杂的法则进行检验和评估。实际上，法则归纳法的任务就是如何结合法则来优化性能。

遗传算法 (genetic algorithm)、神经网络 (neural networks)、递归划分

(recursive partitioning)、核回归(kernel regression)、支持向量机(support-vector machines)以及推进决策树(boosted trees)都是用于数据挖掘尝试的最突出的方法。有关不同的方法和主要的统计理论的精彩论述请参见由哈斯蒂(Hastie)、蒂布什拉尼(Tibshirani)和弗里德曼(Friedman)²⁴合著的*The Elements of Statistical Learning*一书。

客观的技术分析研究

客观的技术人员的工作目的就是发现能够在未来产生盈利的法则。他们的研究方法就是运用能够产生可观察的绩效标准的回溯测试。在这个统计数字的基础上，就可以推断出总体参数、法则的未来绩效或者预期绩效。因此，我们可以说客观技术分析的本身就是统计推断。

为什么客观的技术员必须进行数据挖掘

样本外绩效失败的问题激发了客观技术分析的实践者对数据挖掘的不信任感。这种态度既不明智也不可行。如今，拒绝数据挖掘的客观技术分析师就如同拒绝放弃马车运输的出租车司机一样——这种被废弃的东西已经不再适应让人高效快捷地到达目的地的时候了。

很多因素迫使人们把采用数据挖掘作为获取知识的首选方法。第一，它可以运转。本章后面将要介绍的实验显示了在正常的情况下，进行回溯测试的法则的数量越多，则发现优秀法则的可能性就越大。

第二，技术发展趋势更偏向于数据挖掘。个人电脑的成本效率、利用有力的回溯测试的可能性、数据挖掘软件以及历史数据的有效性都使得个人进行数据挖掘的工作成为可能。而这一切在10年前，受制于成本的数据挖掘只能面向机构投资者。

第三，在数据挖掘演变的现阶段，技术分析缺乏使用传统的科学方法获取知识的理论基础。例如，在物理学这一发达的科学中，单一假设可以从已经创建的理论中推导出来，而它的预测结果也可以用来对新的观察资料进行检验。在缺少理论的情况下，数据挖掘方法提倡对更多的假设(法

则)进行检验。这种鸟枪法^① (shotgun approach) 的风险注定是注定要失败的——法则与数据的相符纯属巧合。我们将在本章的后面部分讨论将这种风险最小化的步骤。

单一法则的回溯测试 VS. 数据挖掘

并不是说所有的回溯测试都可以称为数据挖掘。当单独的法则被提出并且进行回溯测试的时候,就不能把这种行为称为数据挖掘。这种搜索的模式如图 6-2 所示。如果该法则的回溯测试结果不能让人满意,那么就马上停止搜索,并且再考虑更多实际的方法。

回溯测试包括对大量的法则进行回溯测试,并且在测试的最佳绩效的基础上挑选出一种法则。这一过程如图 6-3 所示。注意,这里所显示的是在对第一种法则进行检验之后并没有停止搜索的情况下获得的最初的不能让人满意的绩效。将这种法则进行改善或者重新定义之后,再对其进行检验和绩效的评估。重复这一循环,直到产生具有上佳绩效的法则出现。这一过程包括对成百上千的法则进行检验。

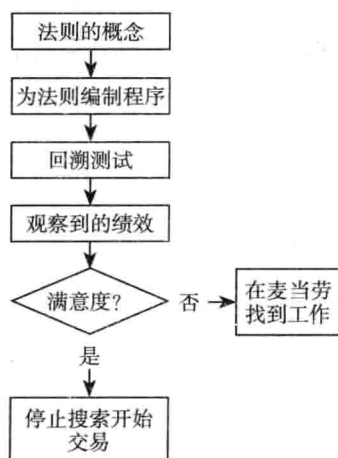


图 6-2 单一法则的回溯测试

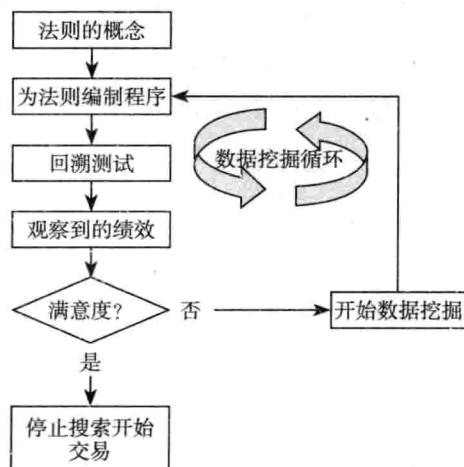


图 6-3 法则的数据挖掘

① 即一种基因组测序方法,将要测序的分子随机切成多个片段进行测序的方法。——译者注

合理运用观察到的绩效

假设绩效统计数据是在回溯测试期间产生的平均收益率。绩效统计数据可以在对单一法则进行的回溯测试和数据挖掘中发挥不同的作用。然而，它们的作用并不相同。在对单一法则进行的回溯测试中，观察到的绩效发挥着对未来法则的估量的作用。而在数据挖掘中，观察到的绩效则发挥着选择标准的作用。对于数据挖掘器来说，当观察到的绩效同时发挥着两种作用，问题就出现了。

在对单一法则进行的回溯测试中，被测试的法则的预期收益率可以合理地用于对法则的预期收益率的无偏估计（unbiased estimate）。换句话说，即进行检验的法则的未来收益率很有可能就是其在回溯测试中产生的收益率。这一点仅仅是我们在第4章中提到的相关内容的重新表述。在那一章节中，我们学习了抽样平均值提供了产生抽样的种群的平均值的无偏估计。抽样平均值也可能由于抽样的随机变化而产生错误。错误的数据可能大于也可能小于总体的平均值，但是没有哪种错误是最有可能发生的。这一原理涉及对单一法则进行的回溯测试。在回溯测试中的法则的平均收益率就是其未来预期收益率的无偏估计。尽管历史平均收益率实际上等于未来的收益率，但是它的绩效可能会更高也可能会更低，但是没有哪种绩效是最有可能发生的（见图6-4）。

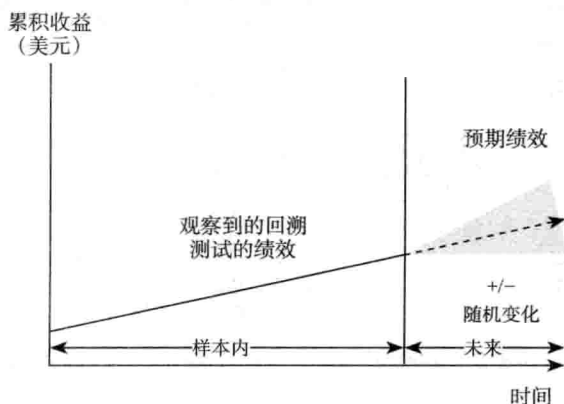


图 6-4 对单一法则进行的回溯测试的预期绩效

在数据挖掘中，回溯测试的绩效统计数据所发挥的作用与其在单一法

则的回溯测试中所发挥的作用截然不同。在数据挖掘中，经过回溯测试的绩效的作用是作为选择的标准。也就是说，它用来确定最佳法则。在所有经过回溯测试的法则的平均收益率当中挑选出产生最高收益率的那个法则。这种方法也是对回溯测试（观察到的）的绩效统计数据的完美合理应用。

从某种意义上说，在回溯测试中产生最高平均收益率的法则实际上就是最有可能在未来有上佳表现的法则，这种说法是合理的。但是这并不能够保证做出最合理的推断。这种表述的最为正规的数学证明请参见由怀特²⁵提供的资料。

数据挖掘器的错误：错误地使用观察到的绩效

我们先来回顾两个重要的知识点。在单一法则的回溯测试中，过去的绩效可以用于对未来绩效的无偏估计。在多法则的回溯测试（如数据挖掘）中，过去的绩效可以作为选择的标准。

数据挖掘器的错误在于，它们使用最佳法则的回溯测试绩效去估计其预期绩效。这种应用是不正确的，因为表现最佳的法则的回溯测试绩效肯定是有偏差的。也就是说，这种绩效水平可以让该法则赢得绩效之间的竞争，并且夸大其预测力和预期绩效。我们将其称为数据挖掘偏差。这一概念如图 6-5 所示。

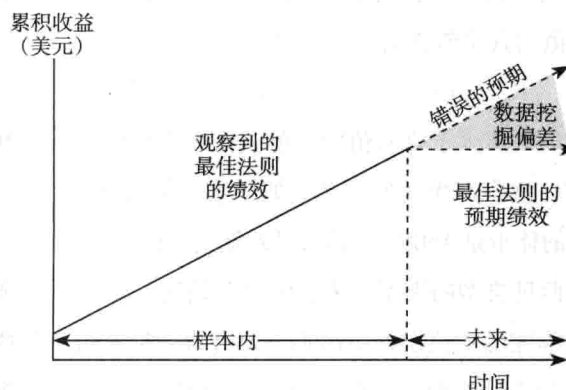


图 6-5 产生最佳数据挖掘法则的预期绩效

最佳法则的绩效不会在样本外真正的下降。它仅仅是通过下面的方式

来表现的,即样本外绩效假设它能够反映出不含运气成分的法则的真实预测力,并帮助其击败其他法则。而样本内绩效则是预测力水平,可能为0,以及含有很大的运气成分的结合体。好运气只在对法则进行回溯测试的过程中起到了作用,但现在则一点作用也没有了。同理,猴子并没有在展示会上丧失它在文学方面的才能。猴子的表现仅仅反映出了当好运气的成分被最小化之后它的真实文学素养,而这种文学素养可以让它碰巧敲出一些能够与莎士比亚诗歌当中的诗句相匹配的段落。

数据挖掘和统计推断

这部分内容讨论的是数据挖掘和统计推断之间的关系问题。主要包括以下几点:①有偏估计量和无偏估计量的区分;②随机误差和系统误差(如,偏差)的区分;③无偏估计量受到随机误差的困扰,而有偏估计量则受到随机误差和系统误差的双重影响;④统计报表适用于大规模的观察资料,比如说大量的估计;⑤数据挖掘偏差显示了许多数据挖掘实例出现的普遍结果,因此,我们不能说任何特定的数据挖掘结果都是具有偏差的。

无偏误差和系统误差

所有科学的观察都会受到误差的困扰。关于误差的定义是这样的:它是观察到的数值与真实值之间的差。即:

$$\text{误差} = \text{观察到的数值} - \text{真实值}$$

当观察到的数值大于真实值时,就会产生正的误差。当观察到的数值小于真实值时,就会产生负的误差。如果某一范围显示,一个人体重140磅,而他实际的体重是150磅,那么误差就是-10磅。

误差包括两种典型的类型:无偏的和有偏的(系统的)。所有的观察资料都受到不同程度的无偏误差的影响。没有哪一种衡量工具或者技术是十全十美的。一种类型的误差的预期值为0。也就是意味着如果在某些现象中使用大量的观察资料,比如说是法则的收益率,那么这些观察资料只会受到无偏误差的影响,对于真实值来说,观察到的值将会被随意分配。如

果计算所有误差的平均值, 那么这个平均值很有可能就是 0 (见图 6-6)。

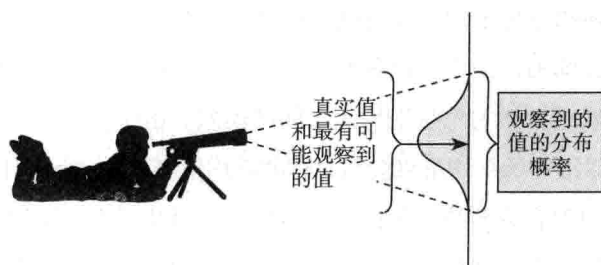


图 6-6 无偏观察

相反地, 受到系统性误差影响的观察资料具有倒向真理某一方面的倾向。我们把这种观察结果称为有偏的。当计算大量有偏观察结果误差的平均值的时候, 这个平均误差将肯定不是 0 (见图 6-7)。图中显示了正偏差的观察结果和正平均误差 (观察到的值 - 真实值)。

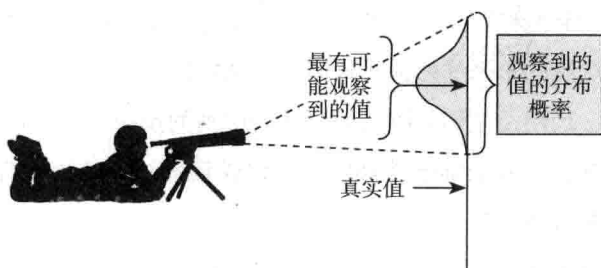


图 6-7 带有系统性误差的观察结果

假设一位化学师对 100 个独立的实验在经过化学反应后留下的残渣的重量进行观察。无偏误差有可能要归结为在对 100 份残渣进行称重的过程中, 实验室湿度的随机变量。残渣湿度的总量影响着其重量。如果存在系统性偏差的话, 就有可能是由于规模的缺陷造成的, 并由此导致观察到的重量低于实际真实的重量。

无偏和有偏的统计数字

对大量的观察结果进行解释是非常困难的。正如我们在第 4 章中所讨论的那样, 首先对数据进行简化才是合理的。我们把大量的测量结果简化为规模更小的概括统计量: 抽样平均值、样本方差以及其他能够描述整个

观察结果的经过计算的计量单位。

第4章还指出像平均值这类的抽样统计量是有可能发生随机误差的一种特殊类型,即抽样误差。这种偏差是无偏的。所以,术语误差是指抽样平均值和提取出抽样的总量平均值之间的误差。由于抽样并不能够完全代表总量,所以抽样的平均值将会与总体的平均值产生一定程度上的背离。

我们再回到客观的技术分析这个话题上,对法则进行的回溯测试产生了大量的观察结果的抽样,即法则的每日、每周以及每月的收益率。这个样本对于把其自身简化成为绩效统计数据(年平均收益率、夏普比率等)来说是非常清楚的。和其他任何的样本统计量一样,绩效统计量会受到随机误差的影响。然而,统计数据也有可能受到系统性误差和有偏误差的影响。

为了表述的更加清楚,历史绩效统计量的作用并不是把钱存入银行或者买一辆法拉利这样简单。它唯一的经济学作用就是根据历史绩效对产生这一数据的法则的未来绩效进行推断。客观的技术人员以置信区间和假设检验的形式,运用经过回溯测试的绩效对法则的预期收益进行推断。不论发生哪种情况,推断的准确性都要以无偏或系统的,以及它们的数量这3种类型的误差为基础。

有偏统计量受到系统误差的困扰。在之前对单一法则进行回溯测试的表述当中,其平均收益率是无偏的统计量。因此,以回溯测试的绩效为基础对法则的预期收益率进行的推断将会导致唯一一种形式的无偏误差,我们将其称为抽样变异性。

然而,通过数据挖掘发现的最佳法则,它的情况却并非如此。观察到的产生最佳绩效的法则的平均收益率是正的有偏统计量。结果,基于此的推断就会产生系统性误差。这意味着当执行一个假设检验的时候,人们对原假设的排斥倾向比显著性水平所显示的内容的频率要高。例如,在显著性为0.05的水平上,人们对排斥原假设误差的预期仅为0.05。然而,如果观察到的绩效是正偏差,那么原假设被排斥的频率就会多于实际的结果,也许会更加频繁。这种情况就会导致进行交易的具有明显预测力的法则在实际中并不拥有预测力的结果出现。但问题是:到底是什么原因导致了经

过回溯测试的最佳法则的收益率，以及由数据挖掘器采用的法则，过分夸大了其实际的预测力呢？

平均值 VS. 最大值

单一法则的回溯测试者和数据挖掘器都在观察两个完全不同的统计量。单一法则的回溯测试者观察的是单一样本的平均值，而数据挖掘器观察的则是大量抽样平均值里的最大平均值。我们非常容易忽略这两个完全不同的统计量的事实。

为了更加明确的说明，在单一法则的回溯测试这一案例中，只存在一组结果，即在回溯测试期中产生的法则的每日收益率。它们是由绩效统计量（例如，每日平均收益率）进行概括的。在数据挖掘中，大量的法则都经过了回溯测试，所以就存在很多组的结果和很多的绩效统计量。如果对 50 种法则进行回溯测试，那么数据挖掘就可以在挑选产生最大平均收益率的法则之前，获得 50 次观察平均收益率的机会。

这组 50 个平均收益率的值可以被看作一组观察结果，相反，它们可以用统计量的方式进行总结。例如，假设要计算平均值的平均值（mean of the means）。也就是说计算所有 50 个法则的平均收益率。而其他对这组观察结果进行计算的统计量则是最小的平均值（minimum mean），也就是表现最差的法则的平均收益率。而另外的一个统计量就是最大的平均值（maximum mean），即表现最佳的法则的平均收益率。

通过数据挖掘观察到的统计量是指在这 50 个法则当中挑选出来的最大平均值。如图 6-8 所示，我们把所有的法则都假设为是无用的，且预期收益率都等于 0。图中每一个圆点代表着观察到的不同法则的平均收益率。请注意，拥有最大收益率（37%）的那个法则的平均收益率并不代表着这个法则的预期收益率（0）。它所获得的收益率只不过是回溯测试中的好运罢了。当对许多法则进行回溯测试时，获得最高平均收益率的那个法则几乎总是受到好运的眷顾。这也就是为什么回溯测试产生的绩效会夸大它的预期绩效。在不同的样本中，那些缺少运气成分的法则，通常表现得更差。

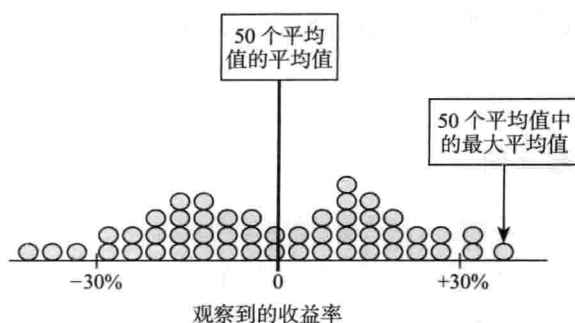


图 6-8 观察到的 50 个法则回溯测试的绩效

要想降低这种情况的发生，必须记住以下几点：在大量的法则当中观察到的具有上佳表现的法则的绩效，是最佳法则的预期绩效的正有偏估计，这是因为它包含了大量的好运成分。对此并不理解的数据挖掘器有可能对这个法则的样本外绩效感到失望。

合理的推断需要正确的抽样分布

合理的统计推断是建立在利用正确的抽样分布的基础之上的。每一个检验统计量都有一个适合于检验其统计显著性的抽样分布。如果观察到的检验统计量实际上是大量样本平均值的最大平均值，那么对检验单一样本平均值显著性做出修正的抽样分布，或者是构建的置信区间就会出现错误。

因此，为了做出合理的推断，数据挖掘器需要获得在大量的平均值中的那个最大平均值的抽样分布，因为在对由数据挖掘发现的最佳法则进行评估的时候，这个抽样分布是需要进行分析的统计量。最大平均值的抽样分布的集中趋势反映了运气成分在数据挖掘中的作用。而单一样本平均值的抽样分布的集中趋势则没有体现出这种作用。

现在我们来分析一下所有的这些因素是如何影响对显著性的检验的。在第 5 章中我们曾经讨论过，在对显著性进行传统的检验时，原假设认为交易法则的预期收益率等于 0 或者小于 0。检验统计量是观察到的法则的平均收益率。在下面的这个案例中，我们假设年收益率是 +10%。检验统计量的抽样分布集中于假设的 0 值这一点。而概率值则指等于或者大于 10% 的收益率的抽样分布的面积。这一面积代表着一种概率，即如果法

则的预期收益率确实等于或者大于0，那么法则就会碰巧得到一个等于或者大于+10%的收益率。如果概率值小于现值，比如说是0.05，那么原假设就会遭到排斥，而当法则的预期收益率大于0时，备择假设（alternative hypothesis）就会被接受。只要一个法则的绩效处于评估中就可以了。

我们现在来分析一下在数据挖掘案例中的显著性检验。我们继续前面介绍的数据挖掘样本，假设对50个法则进行回溯测试。具有最佳绩效的法则获得的年收益率是37%。对于传统的显著性检验来说，在0.05的显著水平上，我们所看到的情况如图6-9所示。图中的抽样分布集中于0，这反映了原假设的内容，即不具有预测力的法则的预期收益率为0。但是这个抽样分布并没有对数据挖掘偏差做出解释。在本章后面进行的实验结果显示，这种假设是错误的。它们证明即使在数据挖掘过程中，所有法则的预期收益率等于0，而具有最佳绩效的法则仍然有可能显示出大于0的绩效。

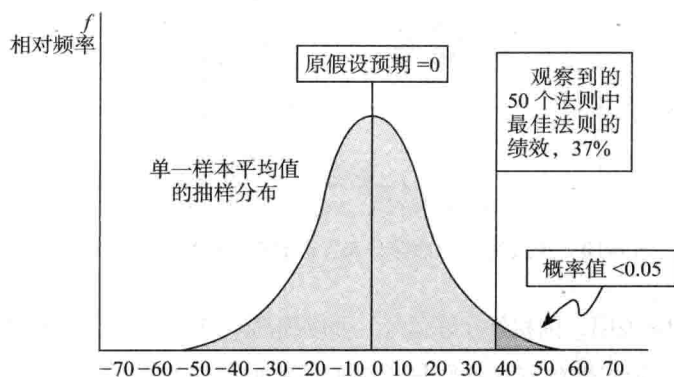


图 6-9 传统的抽样分布（没有考虑数据挖掘偏差）

请注意在图6-9中，在概率值小于0.05的条件下，观察到的+37%的绩效所在的位置是远离抽样分布右边尾部的。在这一实证的基础之上，原假设就会受到排斥，而做出该法则的预期收益率大于0（即法则具有预测力）的结论就是错误的。

然而，如果对50个法则当中那个最佳法则所观察到的绩效进行更加高级的统计显著性的检验，并且把它考虑到数据挖掘偏差中来，则图形所显示的内容就是另一回事了（见图6-10）。图中，对观察到的最佳法则的平均收益率与经过修正的抽样分布进行比较。这就是50个平均值中最大平

均值统计量的抽样分布。这个抽样分布恰当地反映了数据挖掘的偏压效应。我们注意到抽样分布不再以 0 为中心, 而是以 +33% 为中心。根据图 6-10, 最佳法则的绩效不再表现出显著性。抽样分布等于或者大于观察到的收益率 37% 的那部分几乎占据了全部分布面积的一半 (0.45)。换句话说, 如果 50 个法则当中每一个法则的预期收益率等于 0 的话, 那么由于运气的成分所导致的最佳法则的平均收益率大于或者等于 37% 的概率就是 0.45。通过上面的讨论, 我们可以清楚地知道 37% 并不具有统计显著性。特定的数据抽样刚好有利于法则获得突出的表现。

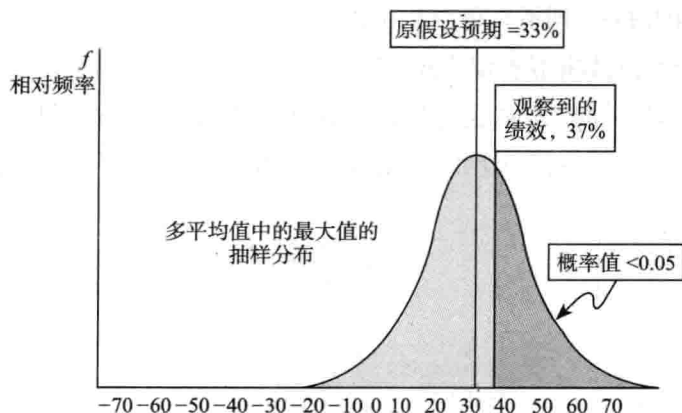


图 6-10 经过修正的抽样分布 (把数据挖掘偏差考虑在内)

图 6-10 表明, 随机性 (好运气) 可以提高不具有预测力的法则的绩效。事实证明, 随机性仅仅是共同导致数据挖掘偏差的两个因素之一。另外的一个因素就是强调多重比较程序的选择原理, 即挑选一个具有最佳绩效的备用法则。

数据挖掘偏差: 两种因素的影响

数据挖掘偏差是受到两种因素共同的影响而导致的: ①随机性; ②数据挖掘的必要选择, 或者任何多重比较程序——挑选一个具有最佳绩效的备用法则。这部分的内容就是要分析这两种因素是如何共同导致观察到的最佳法则的绩效夸大其预期绩效的。

观察到的绩效的两个组成部分

观察到的法则的绩效包含两个组成部分。其一是它的绩效可以归结为法则本身的预测力，如果它有的话。这一绩效的组成要素是由于重复出现的市场行为特征所决定的，而这种特征是可以被法则利用并且在马上实现的预期当中继续证明的。我们将它称为法则的预期绩效。

其二是随机性。随机性可以通过好运气或者坏运气来证明。好运气促使观察到的绩效高于预期绩效，而坏运气则使得观察到的绩效低于预期绩效。随机性并不能够在马上实现的预期当中重复出现。

上面所讨论的内容可以用图 6-11 中的公式来表述。

$$\text{观察到的绩效} = (\text{预期绩效}) +/-(\text{随机性})$$

图 6-11 观察到的绩效的两个组成部分

随机性的范围

对随机性的范围加以考虑是非常有用的。这一概念如图 6-12 所示。²⁶ 在该范围的一端，观察到的绩效由随机性来控制。在这种情况下，所有获得的绩效都纯属运气。在此，我们发现了在文字处理器上跳舞的猴子所创作出来的文学作品，以及买彩票的人的中奖结果。而在该范围的另外一端，则是由理应获得的行为或者系统性的行为来控制。这一点是非常重要的。现在先做出一个有效的数学定理证明。然后在它周围的是观察到的音乐家的演奏水平。在接近演奏水平的地方，我们发现了物理定理，它能够通过某些自然因素的有序行为做出精确的预测。

接近随机范围的中间部分则是最感兴趣的数据挖掘问题。距离范围的随机结束点越远，出现数据挖掘偏差的风险就越大。在靠近连续右边尾部的地方，我们看到一个包括技术分析法则在内的区域。金融市场的复杂性和随机性使我们相信即使是最强有力的法则，它的预测力也是微弱的。在这个区域中，数据挖掘偏差的程度将会被扩大。

现在我们来谈一条重要的原理。在观察到的绩效中，相对于理应获得的结果而言，随机性（运气）的分布越大，则数据挖掘偏差的程度就越大。

原因如下：运气的成分越大，任何一个处于备选法则当中的法则获得异常幸运绩效的机会就越大。而这个法则就是那个被数据挖掘器所挑选出来的法则。然而，在观察到的绩效完全、或者主要是由于备选法则的理应获得的结果造成的情况下，数据挖掘偏差将是不存在的，或者说是出现的概率很小。在这些案例中，备选法则过去的绩效将会是未来绩效的最可靠预估式，而数据挖掘器则很难做到这一点。

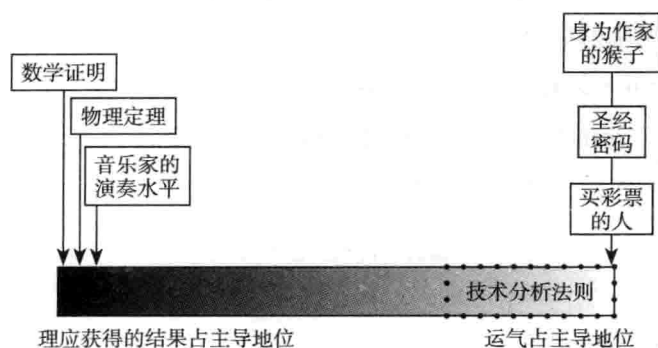


图 6-12 随机性的范围——在观察到的结果中，
理应获得的结果 VS. 运气的成分

由于很难对金融市场进行预测，大多数法则的观察到的绩效都要归结为随机性，而不是其自身的预测力。因此，在图 6-11 中的公式将会在图 6-13 的公式中得到更好的说明。相对于理应获得的结果而言（预测力），随机性的作用越大，技术分析法则的数据挖掘受到更大程度的偏差的影响也就越大。

$$\text{观察到的绩效} = (\text{预测力}) +/-(\text{随机性})$$

图 6-13 在观察到的绩效中，随机性和预测力的相对分布

对应用于多重比较过程的问题的范围，可以把它理解为位于沿着随机性范围的那部分。靠近范围随机的尾部是与技术分析法则的数据挖掘有关的范围。我用这个范围，而不是单独的位置来代表数据挖掘，就是要表明随机性的水平在每一次的挖掘风险中的表现都是不同的。在本章随后将要演示的例子中，存在于任何给定的数据挖掘风险中的随机性水平都是由 5 个独立的因素决定的。²⁷ 当我们把这 5 个因素综合起来进行考

虑的话，就有可能推动能够解决数据挖掘偏差的统计显著性程序的发展。这些程序会减轻运用数据挖掘获取新知的客观的技术分析实践者所面临的主要问题。

在不同的随机性条件下进行多重比较程序的有效性

多重比较程序被用于寻求最佳解决方案。我们提出一个备选解决方案的全域。一个优值（a figure of merit）可以对每一个观察到的法则的绩效进行量化，并且将其中观察到的最佳绩效挑选出来。

很多人认为多重比较程序（MCP）履行了两个承诺：①拥有观察到的最高绩效的那个法则，很有可能在未来表现出最佳绩效；②表现上佳的备选法则的最佳绩效是其未来绩效的最可靠的估值。MCP实现了它的第一个承诺。然而，在对技术分析进行回溯测试的领域，它却没有实现第二个承诺。

关于第一个承诺，即拥有最高绩效的备选法则，很有可能在未来表现出色。这一点在观察资料的数量达到无穷时已经被怀特²⁸证明。怀特认为，当样本容量达到无穷时，拥有最高预期收益率（真正的最佳法则）的备选法则，会认为自己获得所观察到的最佳绩效的概率是1.0。这告诉我们数据挖掘的基础逻辑是正确的！而拥有观察到的最佳绩效的法则就是那个应该被挑选出来的法则。这种假设的有效性将在本章后面的内容中通过数学实验的方式展现出来，即数据挖掘偏差的试验研究。以上得到的结果显示，当有足够的观察资料用于计算法则的绩效统计量时（例如，平均收益率），拥有最高绩效的法则就会获得比随机从进行检验的法则中所提取的法则高得多的预期收益率。为了让数据挖掘成为具有真正价值的研究方法，必须或者最起码也要通过最小有效性的测试。而实际上它确实做到了！

至于第二个承诺，被挑选出来的备选法则的绩效就是其未来绩效的可靠估值。而这也并不是什么新鲜事儿了。詹森（Jensen）和科恩（Cohen）²⁹指出，当MCP应用于在观察到的绩效中发挥重要作用的随机性当中的时候，观察到的具有最佳绩效的备选法则的绩效就会夸大其预

期（未来）绩效。换句话说，当观察到的数据具有显著的运气成分时，它就会获得正偏差估计。在对技术分析法则进行的数据挖掘中出现的就是这种情况。

在随机性不存在，或者因为随机性较低，而不能对挑选出来的备选法则产生任何影响的情况下，观察到的绩效就是最佳备选法则未来绩效的可靠估计。这种情况在音乐比赛或者进行数学证明的过程中比较常见，在实验中，如果不是运气的成分，那么理应得到的结果就是正确的。

从本质上来讲，只要把大量的观察资料用于计算绩效统计量中，由MCP挑选出来的备选法则就有可能是那个在未来表现上佳的法则。然而，在随机性具有显著影响的问题中，观察到的绩效得到了正偏差，而被挑选出来的备选法则的未来绩效很有可能比能够让它赢得绩效竞争的绩效要差。

在低随机性情况下 MCP 的有效性

首先，分析一下把 MCP 应用于具有低随机性的问题中的情况。在这个范围的末端，备选法则的绩效由最优绩效所占据，而 MCP 在识别最优绩效时是有效的。

这个问题我们可以通过为交响乐团聘请一位首席小提琴手的实例来解释。³⁰全部的候选者包括前来应聘的音乐家。他们当中的每一个人都被要求在没有任何排练的前提下，为评委们演奏一首具有挑战性的作品。而评委的评估就是最优值。我们把这种演奏乐器能力的酸性实验称为“视奏”，这种方式是非常有效的，因为任何伟大的演奏都不是偶然出现的。如果把运气的成分当作一种因素，那么成功的概率就太小了。例如，一位伟大的音乐家由于和妻子发生了争吵而感到身体不适，或者在来面试的路上汽车爆胎了，但是这些随机事件对于一位真正的大师级音乐家来说，影响实在是太小了。

在这种情况下，观察到的绩效就是一个精确真实的优值的指标，同时也是未来绩效的最佳预测器。具体内容如图 6-14 所示。每一个候选者的最优值，等于他的预期绩效，我们用图中的箭头来表示。尽管其中一个人的

表现略强一点，但两个音乐家都非常出色。这种围绕在每一位音乐家的优值周围的可能出现的绩效分布过于狭窄了，也就是说，随机性对观察到的绩效的影响是很小的。注意，即使那个更好的备选具有非运气成分的绩效（即在分布的最底端），而稍差的那个备选受到了好运气的眷顾，它们在分布中也没有出现重叠的现象，表现更好的那个备选将会被挑选出来。用另外一种说法表示就是，随机性不会提高那个较差竞争者的表现，而是挑选出具有更高优值的那个候选者。

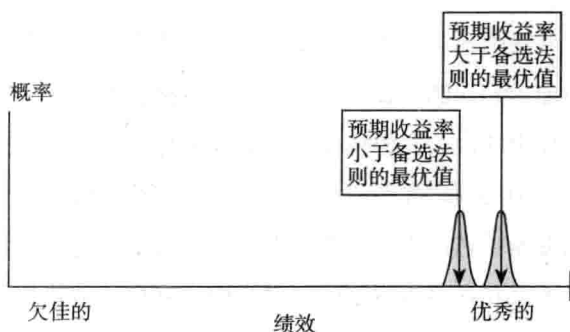


图 6-14 低随机性——最优值的细小偏差。在随机性中显现的真正的最优值

在高随机性情况下 MCP 的有效性

我们现在来分析一下与随机性范围相反的情况，即对技术分析法则的数据挖掘。随机性对被观察到的绩效具有主要的影响，即使是对最有说服力的技术分析法则和拥有相对微弱预测力的模型也是如此。这是对具有复杂性和高度随机性的金融市场的本质做出的推断。对预期收益率为 0 的法则观察到的绩效的概率分布如图 6-15 所示。尽管收益率很有可能为 0，但是由于好运气或者坏运气所导致的收益率有可能更高或者更低。如果某种法则在回溯测试中的运气不佳，观察到的绩效就有可能是负数。对于客观的技术分析人员来说，最麻烦的问题就是当没有任何作用的法则受到了好运气的眷顾，并且得出了正的收益率。这种情况可以让客观的技术分析人员相信他们发现了技术分析的金矿。实际上，法则发出的信号与市场偶然出现的波动是完全相符的。

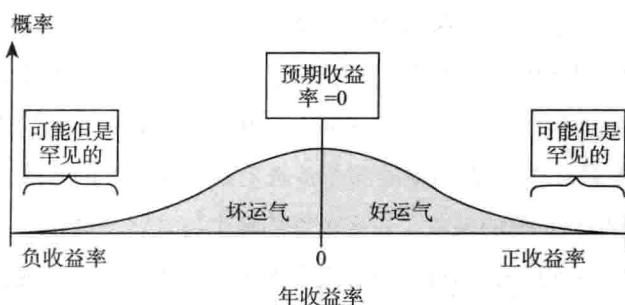


图 6-15 观察到的绩效的概率分布

尽管单一的法则不会出现极端正值或负值的平均收益率，但是随着越来越多的法则加入到回溯测试的行列，这种情况正在增加。正如买彩票的人越多，有人独中两次大奖的机会就会增加一样，对更多的法则进行回溯测试，就会使某种法则增加获得极端幸运的观察绩效的机会。像这种法则就会被数据挖掘器挑选出来。从本质上说：当观察到的绩效主要是由于随机性获得的，那么在众多观察到的绩效当中具有最佳绩效的那个，在很大程度上就是一个随机效应。

为了理解运气成分是如何影响数据挖掘器的，想象一下你正在观察一个数据挖掘器的工作。另外，你可以想象自己非常幸运地知道数据挖掘器所不知道的事情，即每一个进行回溯测试的法则的真实预期收益率。假设数据挖掘器对 12 种不同的法则进行回溯测试，而你知道这 12 种法则的预期收益率为 0。数据挖掘器所知道的事情是通过每一个法则的回溯测试得到的收益率。在图 6-16 中，每一个法则观察到的绩效都由在概率分布上面的箭头表示。每一个分布以 0 为中心，即反映了每一个法则的预期收益率为 0 的事实。请注意有一个法则非常幸运，它所观察到的绩效为 +60%。这个法则就会被数据挖掘器所挑出。在这种情况下，数据挖掘偏差就是 60%。

如果数据挖掘器对更多的法则进行检验，那么获得正观察绩效的概率就会增加。作为我们将要讨论的内容，在数据挖掘中进行检验的法则的数量是影响数据挖掘偏差程度的五个因素之一。在图 6-17 中，我们对 30 个法则进行了检验，其中一个法则产生的绩效达到了 100%。这个法则也会

被数据挖掘器挑选出来。其实你知道该法则的未来绩效会非常糟糕的概率很高。

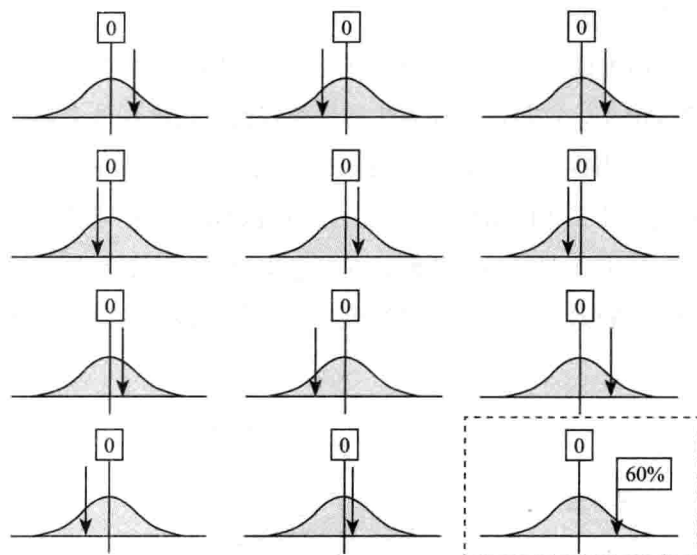


图 6-16 不同的 12 个法则（每一个法则的预期收益率为 0）

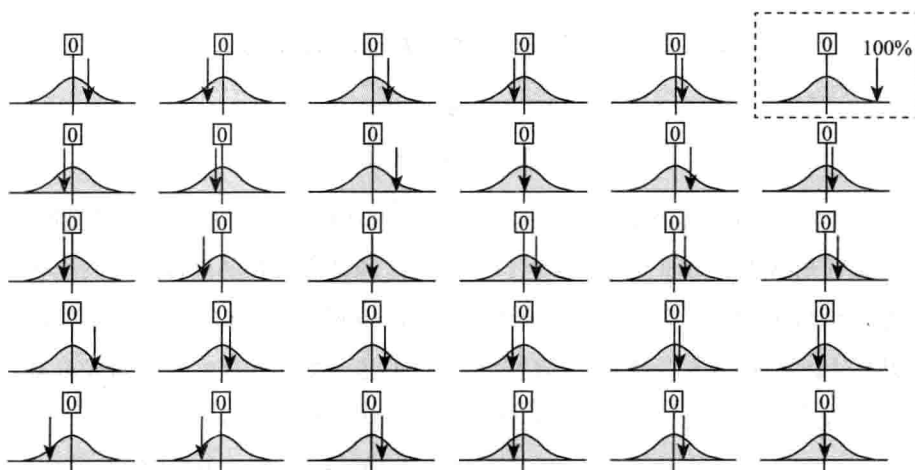


图 6-17 对更多的法则进行数据挖掘，增加了更多极端幸运成分出现的概率

挑选较差法则的风险

之前的例子假设所有的法则都具有相同的值（所有法则的预期收益率

等于0)。当然,数据挖掘器希望在所有经过检验的法则中至少出现一个优秀的法则,而观察到的这个最优法则的绩效也同样会证明它是优秀的。不幸的是,实际情况并非如此。

这就是高度随机性的副作用。具有最高预期收益率的最优法则也许不会被挑选出来,这是因为较差法则所获得的幸运的绩效会赢得数据挖掘竞争。对于出现这种遗憾结果的可能性的直觉(见图6-18)。图中显示了两个具有几乎相同预期绩效的法则的绩效分布,但是其中的一个法则的绩效实际上更加突出。然而,在公布中出现的大量重叠带给人一种感觉,即拥有较差最优值的法则有大量获得最高绩效并且被挑选出来的机会。

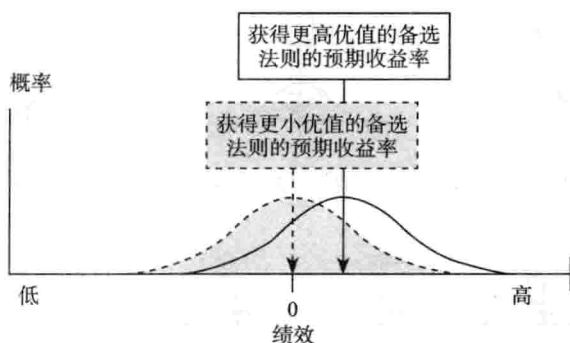


图 6-18 高度随机性和小优值之间的差异

然而,当真实的最佳法则与第二个最佳法则的预期收益率(最优值)的差足够大时,数据挖掘器就可以更加肯定地把观察到的最佳绩效归结于那个绩效更加突出的法则。这种情况如图6-19所示。换句话说,当最佳法则与和它接近的竞争法则的最优值之间的差很大时,优值在这个时候就有可能脱颖而出。这一点可以说误导了数据挖掘器找到金矿的工作。

实际上,数据挖掘器从来就没有注意到挑选较差法则所面临的风险。因为这需要它具备有关真正的优值(预期收益率)方面的知识,而这个值就是我们根本不可能知道的总体参数。就像一首老歌中所唱到的“当未来已成往事”那样,数据挖掘者从来都没有理解数据挖掘偏差的全部结果。

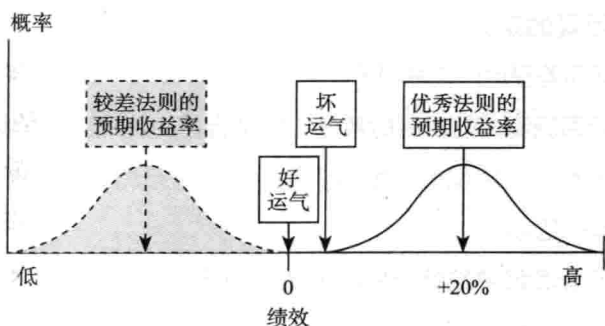


图 6-19 高度随机性和较大的优值差额（优秀的法则战胜了随机性）

决定数据挖掘偏差大小的 5 个因素

我们来简要回顾一下到目前为止所掌握的知识。

- 数据挖掘偏差是指赢得数据挖掘竞争的法则的观察到的绩效与它真实的预期收益率之间的差。
- 观察到的绩效是指从回溯测试的法则中得到的绩效水平。预期绩效是指法则在未来获得的理论绩效。
- 由数据挖掘发现的具有最高绩效的法则的观察到的绩效通常是正偏差的。因此，它的样本外预期绩效将小于能让它战胜其他被检验的法则的样本内绩效。
- 观察到的绩效是随机性和预测力的结合体。随机性的相对分布越大，数据挖掘偏差的程度就越大。

数据挖掘偏差和随机性的相对分布之间的关系如图 6-20 所示。作为将要在下面的章节进行解释的内容，在给定的数据挖掘风险中遇到的随机性的程度，取决于具有风险特征的 5 个因素。

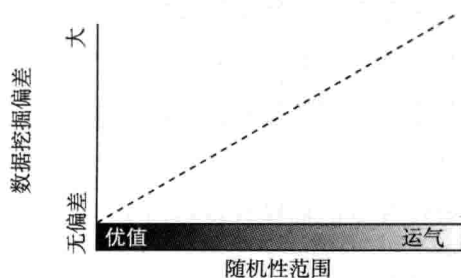


图 6-20 数据挖掘偏差与随机性程度的关系

1.5 种因素的定义

决定数据偏差程度的 5 种因素。

(1) 进行回溯测试的法则的数量。这是指在发现最高绩效法则的数据挖掘过程中进行回溯测试的法则的全部数量。进行测试的法则数量越多, 数据挖掘偏差就越大。

(2) 用于计算绩效统计量的观察资料的数量。观察资料的数量越大, 数据挖掘偏差出现的概率就越小。

(3) 法则收益率之间的关系。这是指进行测试的法则的历史绩效之间的相关度。相关度越低, 数据挖掘偏差的概率就越大。

(4) 异常正收益率值的存在。这是指在法则的历史绩效中出现的比较夸张的收益率, 例如, 在某天出现了比较异常的正收益率。在这种情况下, 尽管当异常正收益率的数量相对于计算绩效统计量的全部观察资料是较小的时候, 这些影响已经降到了比较低的水平, 但是数据挖掘偏差仍然有变大的倾向。换句话说, 更多的观察资料会稀释异常正收益率的偏差影响。

(5) 存在于法则预期收益率中的变差。这是指存在于进行回溯测试的法则中的真正优值(预期收益率)中的变量。变差越小, 数据挖掘偏差出现的概率就越大。换句话说, 当进行测试的法则集合具有相同的预测力程度时, 数据挖掘偏差就会变大。

2. 每一种因素对数据挖掘偏差的影响

图 6-21 ~ 图 6-25 显示了每一种因素与数据挖掘偏差大小之间的关系。

这些关系都涉及特定的数据挖掘风险, 其中正在使用的绩效统计量就是法则的平均收效率。由于像夏普比率的绩效统计量的抽样分布不同于平均收益率的分布, 因此这 5 种因素与数据挖掘偏差之间的关系也是不尽相同的。

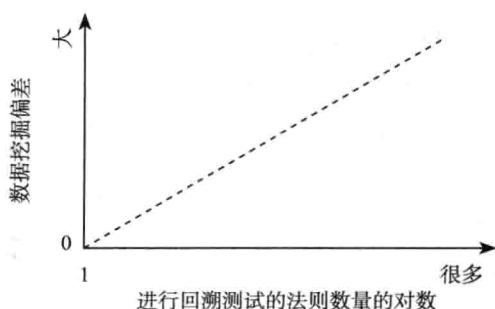


图 6-21 进行回溯测试法则的数量
(第 1 种因素)

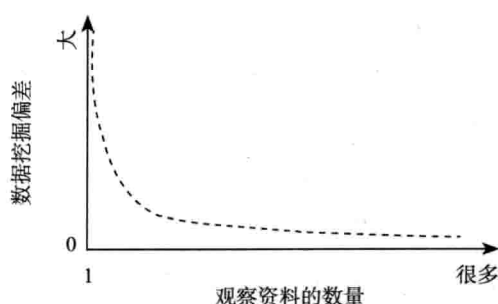


图 6-22 用于计算绩效统计量的观察资料的数量 (第 2 种因素)

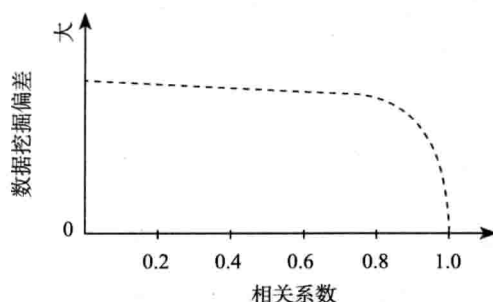
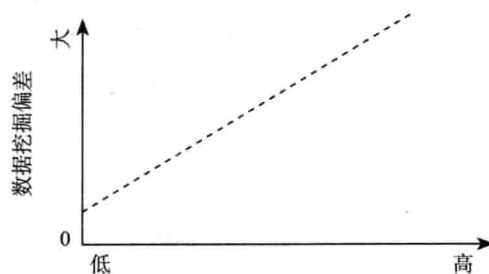


图 6-23 法则收益率之间的关系 (第 3 种因素)



相对于用于计算绩效统计量的观察资料数量的异常值的数量和范围

图 6-24 异常正收益率的存在 (第 4 种因素)

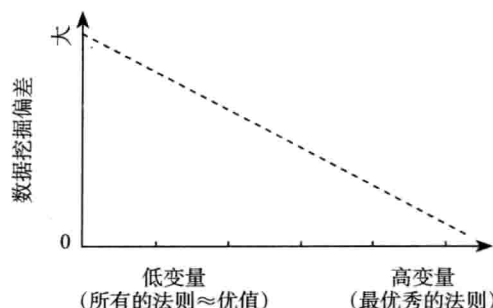


图 6-25 预期 ROI 中的变量——在法则集合中优值的不同程度 (第 5 种因素)

聪明的读者也许会对这种意外感到费解，假设法则的预期绩效是未知的（这里指的是总体参数），而这个参数又是衡量数据挖掘偏差必不可少的一个因素。读者的质疑是正确的；对于实际的法则来说，数据挖掘偏差永远是未知的。然而，对于**虚拟的法则**来说，这是已知的，并且可以进行研究的。虚拟的法则是用于计算机模拟的交易信号，其精确性（正确信号的部分）是可以通过实验进行设定的。这就使我们知晓虚拟的法则的预期收益率成为了可能。前面所述的意外是以一组对虚拟的法则进行的测试为基础的。这些测试的结果我们将在下面的部分进行描述，并且说明 5 种因素对数据挖掘偏差的影响。

数据挖掘偏差的试验研究

科学家根据对观察资料的观测进行推断。有力的推断是建立在精确的观察这一基础之上的。因此，我们最主要的任务就是确定进行观察的程序的准确度。我们通过一种被称为“校准”（calibration）的观察过程来体现衡量随机误差和系统误差的过程。其中一种对观察过程准确程度的校准就是对已知答案是正确的问题进行检验。这种方法使得衡量随机误差和系统误差成为可能。

客观的技术专家就是技术分析领域的科学家。他们最主要的工作就是对法则进行回溯测试。可以观察到的结果被称为绩效统计量。在这个统计量的基础上，可以对法则的预测力或者预期绩效做出推断。因此，客观的技术分析师应当关注在通过回溯测试获得的绩效中出现的随机误差和系统误差。第 4 章说明了由于抽样变异性的原因，绩效统计量有出现随机偏差的倾向。而本章所关注的是来自数据挖掘中的系统性偏差，即数据挖掘偏差。

这部分内容描述的是，通过分析这 5 种因素对偏差大小的影响程度，来研究数据挖掘偏差试验的结果。这一目的可以通过对虚拟的交易法则（ATR）的集合进行数据挖掘来完成。与其他真正的技术分析法则相比，ATR 是非常理想的，这是因为它是受到试验控制的。这就使我们可以用具

有最佳观察绩效的法则来衡量数据挖掘偏差。我们可以知道确切地观察到的绩效，这个绩效统计量对于数据挖掘器来说是已知的，它给出了数据挖掘器希望获得的预期绩效和总体参数。

虚拟的交易法则（ATR）与模拟历史绩效

产生于这些试验的 ATR 历史绩效包括月度收益率。ATR 的预期收益率是通过调整盈利的月度收益率的概率得到控制的。这个概率可以由很多盈利的月份来证明，即大数法则来证明。然而，在包含月份数量较少的样本中，关于特定盈利水平的实际盈利月份所占的比例也许会出现随机性的不同。这个变量把随机性的重要因素引入了试验。而 ATR 的历史绩效就会产生于特定的几个月份当中（如 24 个月）。

在已知可以获得盈利的月份的概率的情况下，以 ATR 为例，它的预期收益率是可以很清楚确定的。相反，我们可以把数据挖掘偏差和将要进行测试的法则结合起来。ATR 的预期收益率可以通过用于计算随机变量的预期值的公式来表示：

预期值 (EV) = p (盈利的月份) $\times$ 平均收益 - p (亏损的月份) $\times$ 平均亏损

由于 ATR 应用于计算从 1928 年 8 月到 2003 年 4 月标准普尔 500 指数的绝对月度百分比变动，所以平均收益与平均亏损的值也是可知的。术语绝对意味着标准普尔指数实际的月度变动信号是被忽略的。在对 ATR 的测试中，由随机过程决定的 ATR 月度收益率的数学符号（+ 或者 -）将在下面进行描述。在这段时间里，包含了近 900 个月份的观察结果，标准普尔 500 指数的平均绝对月度收益率等于 3.97%。因此，计算 ATR 预期收益率的公式等于：

$$\text{预期收益率} = ppm \times 3.97 - (1 - ppm) \times 3.97$$

式中， ppm 是指能够盈利的月份的概率。

ATR 的历史绩效是通过蒙特卡罗模拟法产生的。具体来说，标准普尔 500 指数的绝对月度变动值是从 900 个月的历史数据中重复抽取的。如果特定月份的收益或亏损是确定不变的，那么这个月度收益率就可以通过随机的方式或者计算机模拟的方式获得。这代表着 ATR 当中一个月份的历史

绩效。能够盈利的月份的概率是由实验者所设定的。概率为0.70是指在刻有100个卡槽的轮盘赌中,其中的70个被标记为盈利的,另外的30个则是亏损的。这种产生ATR月度收益率的过程是对特定月份进行的重复,而它的变量则在试验的控制之下。我们用平均月度收益率统计量的方式总结历史绩效,具体过程如图6-26所示。

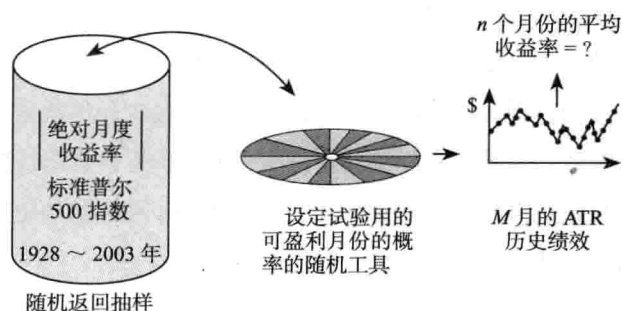


图 6-26 蒙特卡罗 ATR 历史绩效的形成

我们现在来分析一下月度收益概率在4种不同的设定值(1.0、0.63、0.50、0)下ATR预期收益率的表现。把ATR的 p (收益率)设定为1.0(所有可盈利月份),其预期收益率等于标准普尔500指数平均绝对月度收益率(每月3.97%或者年化非复合收益率47.6%)。这个值可以通过插入随机变量的期望公式获得。同样的公式告诉我们月收益概率在设定为0.63的ATR,其预期收益率为12.4%。在0.50的水平上,预期收益率为0。而收益率等于0则获得了每月-3.97%的收益率,或者是年-47.6%的收益率。

请注意,这些预期值是从大量的月份集合中获得的。而在任意数量较少的月份中,观察到的ATR平均收益率都有别于该预期值。包含ATR历史绩效的月份数量越小,预期平均收益率的随机变化就越大。

为了模拟数据挖掘的影响。比如说从10个ATR绩效中挑选出一个最佳绩效,那么就需要产生10个ATR历史绩效。我们选中其中拥有最高平均收益率的ATR,并对其观察到的收益率进行标记。我们将这一程序重复10 000次。这10 000个观察结果将会组成统计量的抽样分布,即从10个ATR绩效中挑选出来的那个最高绩效的平均收益率。

实验 1：对同等价值的 ATR 进行数据挖掘

在试验的第 1 步，设定所有的 ATR 都具有相同的预测力。我们把所有 ATR 的 p （收益率）设定为 0.50，因此，它们的预期收益率等于 0。

1. 因素 1：进行测试的法则的数量

在其他条件不变的情况下，通过回溯测试发现最佳法则的法则数量越多，数据挖掘的偏差也就越大。在键盘上跳舞的猴子越多，敲出的字母与诗歌相匹配的幸运概率越高。同理，进行回溯测试的法则数量越大，则获得超常运气的概率也就越大。

在下面的测试中，每一个 ATR 都对 24 个月这一时间段进行了模拟。首先，来看一下非数据挖掘的情况，即对单一法则进行测试。这里面是不存在数据挖掘和数据挖掘偏差的。尽管 ATR 的 p （收益率）设定为 0.50 意味着预期收益率等于 0，但是抽样变异性使得任何 24 个月的历史绩效产生大于 0 或者小于 0 的平均收益率成为可能。

为了表示 1 000 个 ATR 的历史绩效，在 24 个月当中每一个月的绩效都是通过电脑模拟的。图 6-27 显示了平均年化收益率这一统计量的抽样分布。和我们预期的一样，分布集中于 0，实际上，ATR 的预期收益率的 p （收益率）值设定为 0.50。同理，平均收益率显示了围绕在 0 周围的广泛的变量排列，因为 24 个观察结果是相对较小的样本容量。少量的 ATR，即在抽样分布右侧尾部之外的部分，幸运的成分是相当高的（可以盈利的月份超过 50）。在分布左侧的尾部，我们发现了运气不佳的历史绩效。然而，抽样分布的中心可以对此进行解释。平均来说，在非数据挖掘的条件下，没有预测力的 ATR 的预期平均收益率为 0。在下面的测试中，当最佳 ATR 绩效从两个或更多的 ATR 中选出的时候，拥有最佳绩效的 ATR 很有可能获得大于 0 的收益率，即使它的预期收益率等于 0。

现在我们对从两个或者更多的 ATR 中选出的最佳 ATR 进行数据挖掘。为了分析数据挖掘偏差大小与被选中的最佳法则数量之间的关系，我们要对 ATR 集合的容量进行调整，即被选中的最佳法则的数量。例如，如果 ATR 的数量设定为 10 个，那么从这 10 个法则中选出一个最佳绩效，并记

录下观察到的平均收益率。重复这一过程 10 000 次，并绘制统计量（观察到的最佳 ATR 的收益率）的抽样分布。全部的历史绩效为 24 个月，并把所有法则的 p （收益率）值设定为 0.50（预期收益率等于 0），且所有法则的收益率都是独立的（不相关的）。由于所有 ATR 的预期收益率等于 0，所以数据挖掘偏差只是等于在 10 000 次重复当中观察到的最佳绩效的平均收益率。

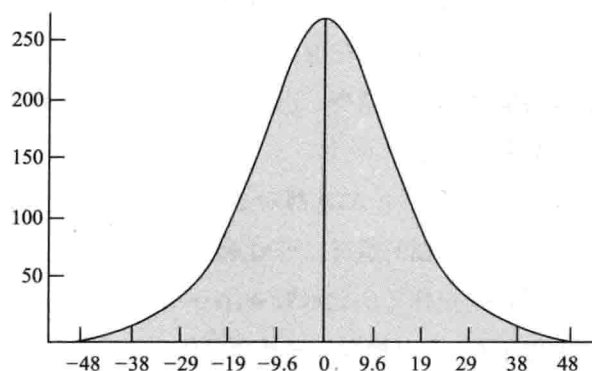


图 6-27 平均年化收益率的百分比
(单一 ATR、24 个月的历史绩效、1 000 次重复)

图 6-28 显示了两个 ATR 是获得最佳绩效的那个法则的抽样分布。抽样分布集中在 8.5% 这个值上。换句话说，在只有两个 ATR 进行测试，且最佳绩效的法则被选中的情况下，数据挖掘偏差是 8.5%。因此，数据挖掘器对于未来绩效的预期是 8.5%，但是，我们知道其真实的预期收益率是 0。

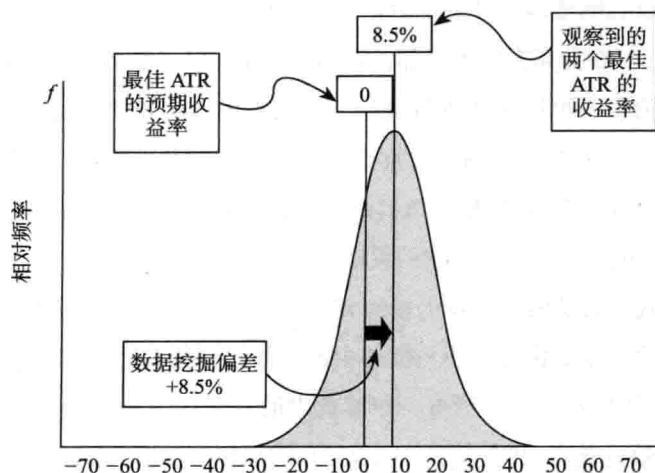


图 6-28 抽样分布——两个最佳 ATR 的平均收益率

如果把 ATR 的数量增加到 10 个，并且从中选出一个最佳绩效，那么数据挖掘偏差就会增加到 22%。然而，我们已知被选中的 ATR 的 p （收益率）值为 0.50%，且收益率为 0。图 6-29 与图 6-28 类似，只不过它所显示的是所观察到的 10 个最佳 ATR 的平均收益的抽样分布。其数据挖掘偏差值为 22%。

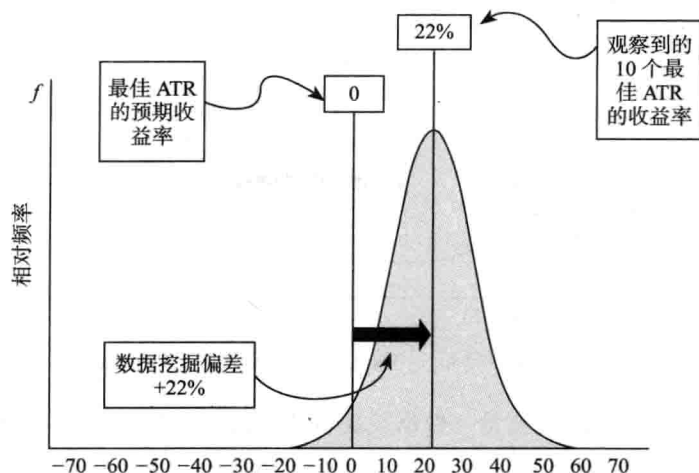


图 6-29 抽样分布——10 个最佳 ATR 的平均收益率

当数据挖掘的集合规模扩展到 50 个 ATR 的时候，数据挖掘偏差则上升到 33%。图 6-30 显示了 50 个最佳 ATR 的平均收益率的抽样分布。

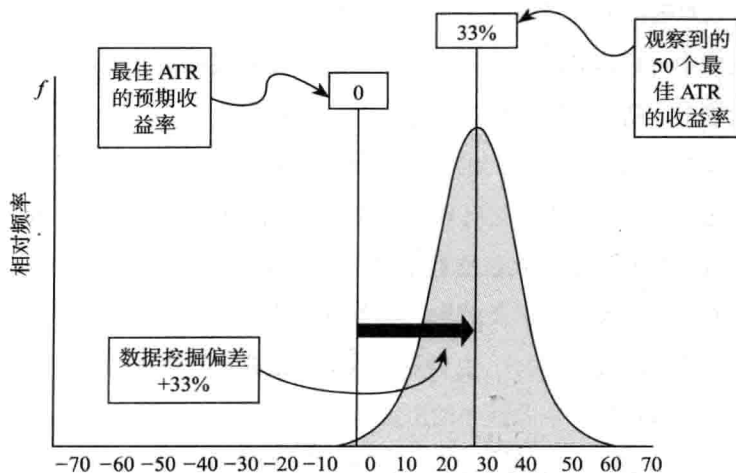


图 6-30 抽样分布——50 个最佳 ATR 的平均收益率

同理，图 6-31 显示了当 ATR 数量增加到 400 个时，其数据挖掘偏差上升到了 48%。

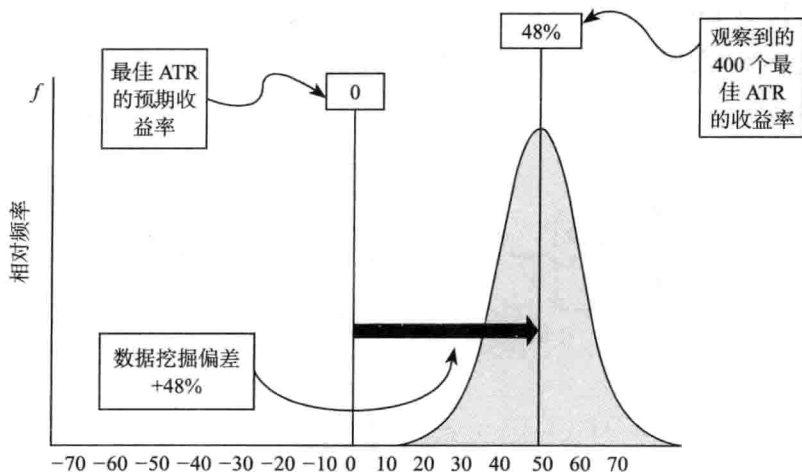


图 6-31 抽样分布——400 个最佳 ATR 的平均收益率

进行测试的 ATR 数量与数据挖掘偏差之间的关系如图 6-32 所示。图中概括了所得到的最佳 ATR（2、10、50、400）的结果。纵轴代表数据挖掘偏差的程度，即观察到的拥有最佳绩效 ATR 的收益率与其预期收益率之间的差。横轴代表了通过回溯测试发现最佳法则绩效的 ATR 数量。它们之间的关系由图 6-32 中的 4 个数据点连接而成的曲线概括。当进行测试的法则数量以对数的形式显示时，即对数（基数为 10），这条曲线呈线性状。从这些试验中可以得到以下的几个要点：为了发现最佳法则而进行测试的法则数量越多，则数据挖掘偏差就越大。

需要指出的是，在之前的测试中显示的数据挖掘偏差的特定量值仅仅在以下这项特殊实验中才是有效的：我们设定标准普尔 500 指数 24 个月的月度绩效，所有法则的收益率都是独立的，且所有法则的预期收益率都等于 0。在不同的数据挖掘尝试中，随着设定值的不同，即使应用同样的原理（法则越多，偏差越大），也会导致数据挖掘偏差的特定水平出现变化。

例如，我们把平均收益率的历史绩效的月份设定为更长的 48 个月，取代之前所用的 24 个月，ATR 的平均收益率的分布就会更加紧密地集中在预期收益率为 0 的这个数值周围。这仅仅是大数法则的一种表现。结果，拥

有最高绩效的 ATR 产生的偏差就会减少。以上内容告诉我们用于计算绩效统计量的观察资料的数量是决定数据挖掘偏差程度的重要因素。

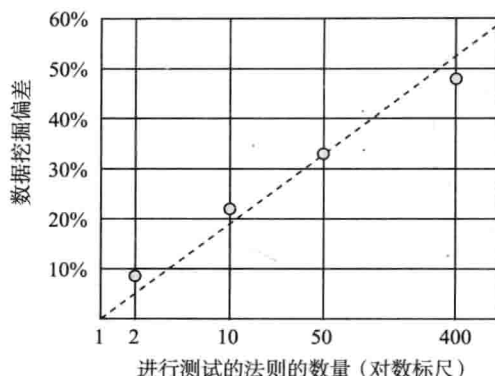


图 6-32 数据挖掘偏差 VS. 进行测试的法则的数量 (年化收益率):
24 个月的历史绩效

我们来分析一下以前的内容。用于计算绩效统计量的月度观察资料的数量越多，统计量的抽样分布的离差就越小。换句话说，用于计算绩效统计量的样本容量越大，出现在观察到的绩效中的随机性和产生极端幸运绩效的概率就越小。不论观察到的绩效中的随机性程度降低到什么程度，都会降低数据挖掘偏差。

数据挖掘偏差的样本容量的重要性如图 6-33 所示。图 6-33 与图 6-32 类似，纵轴都是对数据挖掘偏差的描述，而横轴表示的是相对于发现最佳法则的 ATR 数量的函数。然而，在图中我们用 4 条曲线代替了之前所用的单一曲线。每一条曲线都是以计算每一个法则的平均收益率的月度观察资料的不同数值为基础的，即 10 个月、24 个月、100 个月以及 1 000 个月。以小数点表示的是 24 个月的曲线，其与图 6-32 所示的内容相一致，这两个图是值得我们注意的。首先，所有的曲线随着进行回溯测试的法则数量的增加而上升。这与我们之前对所发现的结果进行讨论的部分相一致，即数据挖掘偏差随着进行测试的 ATR 的数量的增加而增加。其次，也是最重要的一点，就是数据挖掘偏差的程度随着用于计算平均收益率的月份数量的增加而减少。例如，当只有 10 个月的数据可以使用时，那么最佳的 1 024 个 ATR 的数据挖掘偏差大约是 84%。当我们用 1 000 个月的数据来

计算,则数据挖掘偏差就会减少至12%。在接下来的内容中,我们要了解为什么观察资料的数量对于决定数据挖掘偏差的大小是如此的重要。

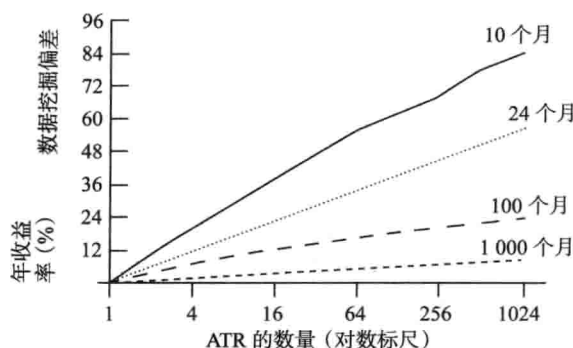


图 6-33 数据挖掘偏差 VS. 用于测试不同样本容量的 ATR 数量

2. 因素 2: 用于计算绩效统计量的观察资料的数量

图 6-33 告诉我们随着用于计算绩效统计量的观察资料的数量的增加,数据挖掘偏差的程度就会减少。实际上,在影响数据挖掘偏差的所有因素中,样本容量是最重要的。使用的观察资料越多,则那些导致得出高平均收益率的幸运的观察结果(盈利的月份)的概率就越小。我们在第 4 章中了解到,当用于计算绩效统计量的观察资料的数量很大时,抽样分布的宽度就会缩小。当观察资料的数量较少时,样本的平均值就会大大地高于或者低于其实际的总体平均值。从广阔的抽样分布中总结出的信息告诉我们,在回溯测试中由于幸运的成分产生出高盈利收益率的概率比靠真实预测力获得的收益率的概率要高。图 6-34 显示了同一法则的两种抽样分布,其预期收益率为 0。请注意,基于少量观察资料的抽样分布更为宽广。从时间较短的历史绩效中计算得出的平均收益率要比从大量的观察资料(更长的历史绩效)中计算得出的平均值更容易得出具有幸运成分的结果。

实验的下一步要分析的是用于计算 ATR 平均收益率的月度观察资料的数量对于数据挖掘偏差的影响。月份的数量不是按照从 1 ~ 1 000 的顺序进行排列,而是以每 50 个月的数量增加,并且计算出每组增量的数据挖掘偏差。我们将其分为两种情况:最好的 10 个 ATR 和最好的 100 个 ATR。和以前的实验一样,所有 ATR 的预期收益率被设定为 0。

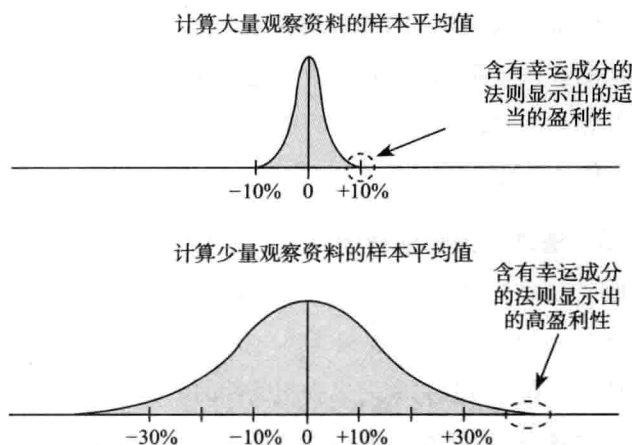


图 6-34 狭窄的抽样分布 VS. 广阔的抽样分布

图 6-35 显示了纵轴上的数据挖掘偏差与在横轴上用于计算 ATR 平均收益率的观察资料数量之间的关系。因为所有 ATR 的预期收益率都等于 0，其数据挖掘偏差也等于观察到的最佳 ATR 的平均绩效。因此，纵轴就被称为数据挖掘偏差，也可以称为最佳法则的平均绩效。注意，随着用于计算平均收益率的观察资料数量的增加，数据挖掘偏差的倾斜度会随之下降。这就是大数法则在数据挖掘中起到的作用。

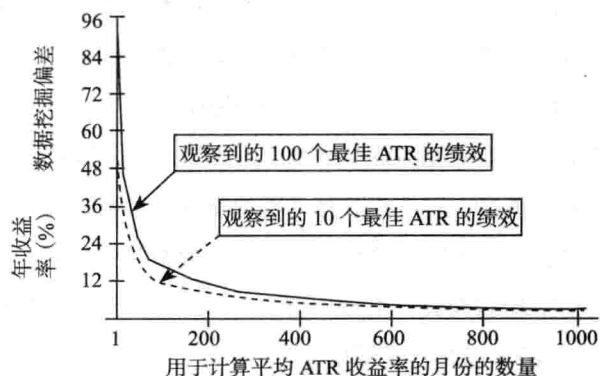


图 6-35 数据挖掘偏差 VS. 观察资料的数量

我们在此处获得的经验是：对从少数信号或者时间区间中计算得出的绩效统计量提出质疑。从图 6-35 中获得的另外一个重要信息是，当观察资料的数量（月份）足够大时，比如说是 600 或者更多，那么在 10 个最佳

ATR 与 100 个最佳 ATR 之间的数据挖掘偏差将会变得很小。这一发现对于数据挖掘器来说是非常重要的推断, 即当观察资料的数量十分充足时, 我们可以在没有显著增加数据挖掘偏差的情况下, 对更多的法则进行数据挖掘。这就是大数法则的规则。

3. 因素 3: 法则相关性的程度

影响数据挖掘偏差大小的第 3 个因素是在众多进行测试的法则之间的相似程度。当产生的历史绩效高度相关时, 我们就称其为相似。也就是说, 它们产生的月度或者每日收益率是具有相关性的。受到测试的法则之间的相关性越强, 数据挖掘偏差的程度也就越小。相反, 法则收益率之间的独立统计程度越强, 数据挖掘偏差就越大。

这种说法之所以合乎情理, 是因为法则之间的相关性的增加导致了影响进行回溯测试的法则的数量的减少。想象下一组法则是完全一样的情况。实际上, 它们会产生完美的相关历史绩效。事实上, 这些法则的集合实际上就是一个法则, 我们已经知道当只有一个法则进行回溯测试的时候, 数据挖掘偏差就会降为 0。我们现在回到这个极端的例子中来, 当法则高度相似时, 就会产生高度相关的收益率, 那么获得超额幸运绩效的概率就会降低。法则之间的相似度越小, 某一法则获得与历史数据相同, 并且获得更高绩效的概率就越大。所以, 法则间的高度相关性降低了影响进行回溯测试的法则的数量, 同时也降低了数据挖掘偏差。

在实验期间, 当对数据挖掘中特定的法则形式的参数持乐观态度时, 法则的相关性就会很高。假设对双移动平均交叉法则持乐观态度, 这意味着每一个受到测试的法则除了参数值之外, 其他都是一样的, 即用于计算短期和长期移动平均的天数。用 26 天和 55 天计算的法则, 其收益率与用 27 天和 55 天计算的法则具有高度相关性。

图 6-36 显示了最佳法则(纵轴)的数据挖掘偏差与法则相关性程度(横轴)之间的关系。我们对法则相关性的模拟如下: ATR 初始的历史绩效是已经产生的。在产生第 2 个 ATR 历史绩效的过程中, 用带有偏差的随机策略来确定第 2 个 ATR 的月度收益率。我们把这个偏差设定在理想的相关度水平上。例如, 如果期望的相关度为 0.7, 那么模拟掷硬币时出现正面

(头像)的概率就等于0.70。如果硬币落地后出现的是正面,那么前两次的ATR月度收益率就是相同的。重复这一行为,并得出下面ATR的历史绩效。

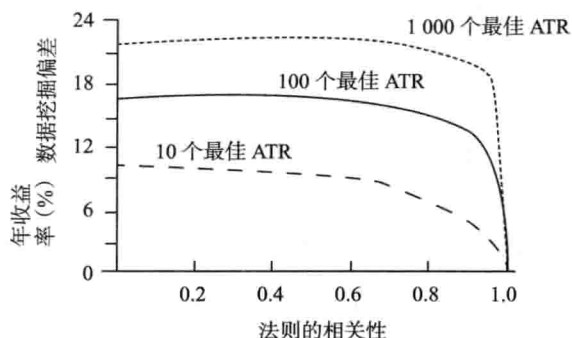


图 6-36 数据挖掘偏差 VS. 法则的相关性
(月度观察资料数量的年化收益率百分比=100)

每一个 ATR 都是用 100 个月的历史绩效进行模拟的。我们进行 3 组不同的测试,每一组基于不同的 ATR 数量:10、100、1 000。换句话说,一组是对 100 个法则中最佳的那个 ATR 进行数据挖掘偏差的衡量。和之前进行的测试一样,所有法则的预期收益率都等于 0。图中纵轴代表了数据挖掘偏差。我们在此处进行分析的因素、法则的相关性,是按照从 0~1.0 的顺序进行改变的。请注意,在法则的相关性达到 1.0 以前,数据偏差仍然会处于更高的水平。因此,在出现收益率高度相关性之前,法则的对于降低数据挖掘偏差并没有起到重要的影响。同样,我们还注意到 1 000 个最佳法则的数据挖掘偏差要比 10 个最佳法则的数据挖掘偏差高。这与因素 1 是相符合的,从测试中挑选的最佳法则的数量越多,数据挖掘偏差就越大。

4. 因素 4: 法则中正异常值回报率的存在

包含少数极端正观察结果的法则收益率样本,具有产生较大数据挖掘偏差的可能性。数据挖掘偏差的大小取决于这个数值的极端程度和用于计算平均值的观察资料的数量。如果使用的观察资料数量很少,那么产生的极端正值会增加样本的平均值。使用大量的观察资料会减少偶然出现的极

端值所带来的影响。

包含极端值的分布说明其具有厚尾性。分布的尾部是指在左边和右边突出的区域,那些极端但是稀少的观察结果就分布在这片区域中。厚尾分布距离分布的中心要比轻尾分布距离分布的中心远得多(见图 6-37)。图 6-37 中最上面的部分呈正常的钟形形态。注意观察其中一条线远离分布中心的速度。这说明极端的观察结果对于实际上根本不存在的观点来说是非常稀少的。人类的身高也是呈正态分布的。我们很难看到身高超过 7 英尺的人,也从来没有人见到过身高超过 10 英尺的人。图 6-37 中下面的部分就属于厚尾。尽管出现极端的观察结果比较少见,但是它们的出现频率要比轻尾分布更大。如果人类的身高也出现这种分布的话,你就会见到身高超过 20 英尺的人了!

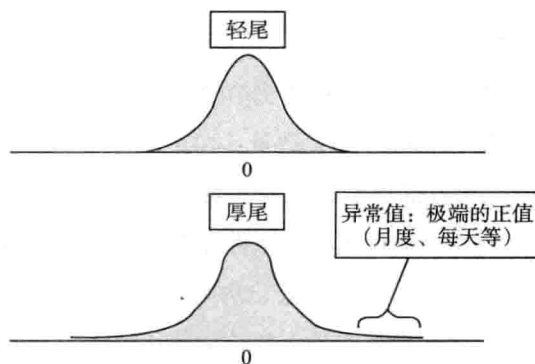


图 6-37 轻尾 VS. 厚尾收益率分布

正如之前所提到的那样,当每日(每周、每月)收益率的分布包含极端的观察结果、单一的平均值的时候,那么从这些观察资料中计算得出的样本平均值也会呈现出极端值。当一个或者多个极端值出现在同一样本中,或者说这个样本容量非常小的时候,就会出现这种情况。出现在收益率分布中的厚尾现象具有导致平均值的抽样分布出现尾重的倾向。然而,样本平均值却从来没有发生过在大多数样本中出现的极端值那样的情况。这就是平均值的影响。因此,相对于区间收益率的分布,平均收益率的抽样分布通常会出现轻尾(见图 6-38)。图中描述了两个收益率分布。左边所示的是带有厚尾的收益率分布,平均值的抽样分布是以同样具有厚尾的分布为

基础的。右边所示的是轻尾的收益率分布。平均值的抽样分布是以其轻尾为基础的。

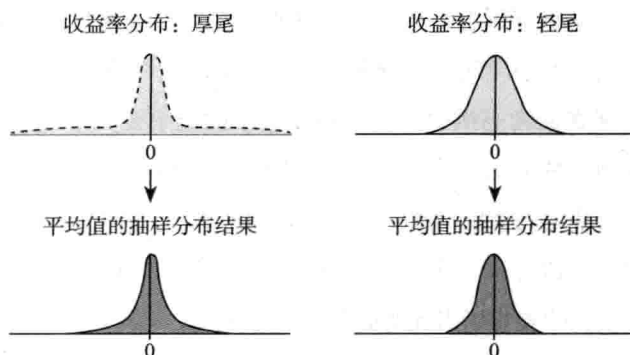


图 6-38 收益率分布的极端值对于抽样分布的影响：同等的样本容量

注意，图 6-38 中的两组抽样分布都是以同样的样本容量为基础的。样本容量是影响抽样分布尾重的一个独立因素。这里保持样本容量不变，以免造成混淆。

现在我们来分析样本容量对抽样分布的尾重造成的影响。当用于计算样本平均值的观察资料越来越多的时候，对收益率分布中的尾重的影响就越小。用于计算样本平均值的观察资料的数量越大，抽样分布的尾部就越轻（见图 6-39）。图中，收益率的分布明显具有厚尾。左边的平均值的抽样分布是以小样本容量为基础的。右边的则是来源于同样的收益率分布，只是用于计算样本平均值的是较大的样本容量。这个抽样分布的尾部是轻尾。

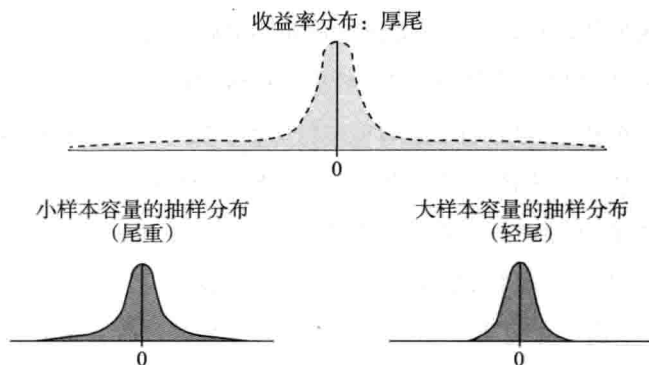


图 6-39 样本容量对抽样分布尾重的影响

上面的内容与数据挖掘偏差又有怎样的联系呢? 带有厚尾的抽样分布是指随机性在平均收益率中起到了重要的作用。对于数据挖掘者来说,³¹ 随机性可不是什么好消息。假设两个法则, 它们的预期收益率都为 0。其中的一个法则具有厚尾的抽样分布。另外一个的抽样分布则是轻尾。厚尾告诉我们法则在回溯测试中产生平均收益率的概率要大于其预期收益率。如果这个法则被数据挖掘器挑中, 其数据挖掘偏差将会很大, 并且样本外绩效也会得出让人失望的结果。在图 6-40 中, 厚尾抽样分布显示了法则在回溯测试中获得超过 30% 的平均收益率的概率范围。在抽样分布具有轻尾的情况下, 较高的收益率出现的概率几乎为 0。

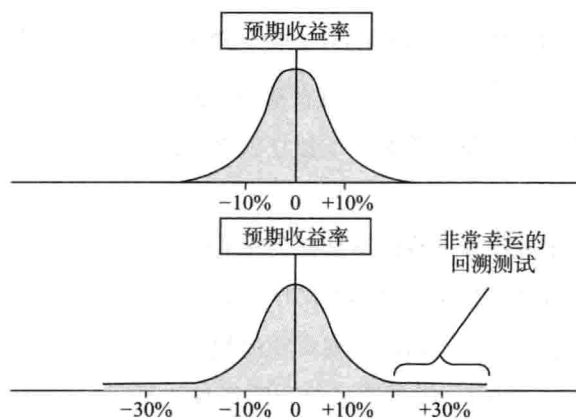


图 6-40 轻尾 VS. 厚尾抽样分布

众所周知, 与正态分布³² 相比较, 股票市场和其他金融资产的收益率分布都会相对出现厚尾。一个持有看空观点的法则碰巧在市场崩盘期间赢得了数据挖掘绩效的竞争, 但是其胜出的理由则是错误的。

尾重的情况和其导致的结果对数据挖掘偏差的影响, 我们将通过实验的方式进行分析。我们对两组 ATR 进行比较。其中一组使用标准普尔 500 指数的月度收益率来产生历史绩效, 其尾部的分布较正态分布要重。另外的一组 ATR 的收益率分布则有意使用轻尾。正如在其他的实验中那样, 观察到的那个具有最佳绩效的 ATR 被选中, 并记录下了它的数据挖掘偏差。这些实验的结果一目了然。数据挖掘偏差对于以轻尾收益率分布为基础的 ATR 来说并不显著。只有当进行数据挖掘的法则存在于固定的目标当中,

这一结果才具有实际的价值。它具有消除极端结果的效果。而这一结果在降低绩效统计量的抽样分布中的尾重方面也发挥了作用。然而，研究者则继续对由于金融市场行为产生的收益率分布进行研究。由于市场极易遭受极端事件的冲击，法则的平均收益率也有获得极端值的趋势。

5. 因素 5：存在于 ATR 预期收益率中的变差

因素 5 是指在数据挖掘过程中，存在于法则集合中的预期收益率的变差程度。我们到目前为止讨论过的 ATR 实验主要考虑的是具有一致预测力的 ATR。它们所有的预期收益率都设定为 0。当所有法则都具有同样的优值，就不会出现真正的优值，而在所观察到的绩效中出现的或低、或高，以及其他的差异则完全是由于运气所导致的。在这种情况下，在数据挖掘绩效竞争中脱颖而出的那个法则就是最幸运的法则。因此，观察到的平均收益率与其预期收益率的差就会很大。所以，当所有的法则拥有相同的预期收益率时，数据挖掘偏差将会扩大。

然而，如果在数据挖掘过程中对全部法则进行的研究不同于其预期收益率的话，那么数据挖掘偏差就会缩小。在某种情况下，数据挖掘偏差会非常小。下面的实验要研究的是预期收益率与全部法则之间的变差对数据挖掘偏差的影响。我们通过建立一个含有不同预期收益率水平的 ATR 的集合来完成实验。回想一下，ATR 的预期收益率完全是由经过设定的盈利月份的概率来决定的。

我们首先分析下面的实验。假设一个 ATR 的全域包含一个真正超过其他 ATR 的法则。该法则的预期收益率是 20%，而其他法则的预期收益率则为 0。我们希望通过在大多数时间中所观察到的最佳绩效来体现出该法则的优势（在其他测试中则不允许这样做）。结果，在大多数时间里，这个优秀的法则都被挑选了出来，而与其相对应的数据挖掘偏差也趋于缩小。数据挖掘偏差之所以会缩小，是因为其获得的绩效是通过本身的预测力，而不是运气的成分获得的，只是运用的方法比较老式而已。

实验 2：用不同的预期收益率对 ATR 进行数据挖掘

下面的测试研究的是在由不同预期收益率的 ATR 组成的全域中的数据

挖掘偏差。在这些测试中，大多数 ATR 的预期收益率等于 0 或者接近 0。然而，还有一些收益率显著大于 0 的 ATR。这符合了对法则进行数据挖掘的人希望遇到的那个可以复制的世界，即在数据的矿山中确实能够找到某些闪光的技术分析。

除此之外，我们还要研究数据挖掘偏差的程度，这组测试又解决了另外一个问题：数据挖掘是以合理的假设为基础的吗？也就是说，挑选出的那个具有最佳绩效的法则在识别真正的预测力时要比从全域中随机抽取的法则做得更好吗？数据挖掘必须，最起码也要通过最低限度的有效性测试，来证明它是一种合理的研究方法。图 6-41 显示了已经给出的 ATR 的指定盈利月份的概率。请记住这幅曲线图特指用于模拟实验的指定月度收益率的概率。

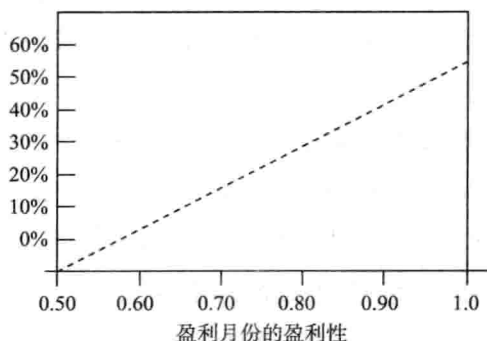


图 6-41 盈利月份的盈利性的预期收益率

回顾一下 ATR 的预期收益率是由其能够盈利的月份的概率（PPM）所决定的。下面的测试假设具有高预期收益率的法则是非常罕见的。为了对这一假设进行模拟，我们将分配给每一个 ATR 的 PPM 从 PPM 的分布中随机提取。用于这一目的的分布与社会上财富的分布图极为相似。我们看到在分布的右侧之外出现了一条长尾，这表明像比尔·盖茨这种富有的人所占的社会财富是一个极端的比率，而其他大部分的人都非常的贫穷。用于我们实验当中的预期收益率的分布被定义为 ATR 的预期收益率为每年 19%，或者说出现这种结果的次数是 1/10 000。在全域中，ATR 平均的预期收益率是每年 1.4%，其大致的分布图形如图 6-42 所示。

尽管构成这一分布的特定假设并不是很精确，但是在一个数量级的范围内也许就是正确的。因为它是以运用技术分析方法的最佳对冲基金的性能为基础的。然而，为了进行实验，特殊性并不是主要的。我们真正的目的是衡量在全域中预期收益率并不相同的法则的数据挖掘偏差，并且研究这一偏差是如何阻碍数据挖掘的客观性的，从而发现在全域中本来就非常

罕见的优值。

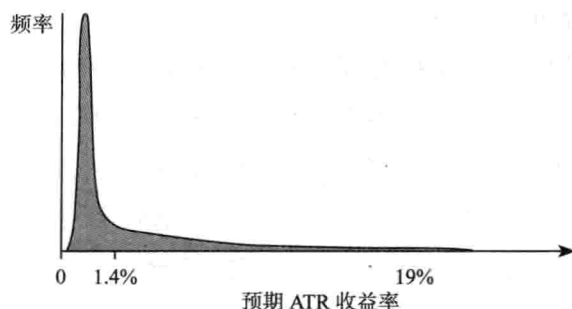


图 6-42 ATR 预期收益率的分布

数据挖掘是以合理的假设为基础的吗？——是的

这部分内容要讨论的是一个重要的基础性问题：数据挖掘是以合理的假设为基础的吗？这实际上是两个问题。首先，在观察到的最高平均收益率的法则比随机抽取的法则具有更高预期收益率的情况下，所观察到的平均收益率是衡量法则优值的最有效指标吗？如果随机抽取的法则在未来的表现与在回溯测试中法则的表现类似的话，那么数据挖掘就是毫无意义的行为。其次，数据挖掘越多越好吗？也就是说，对大量的法则进行测试就真的能够发现具有更高预期收益率的法则吗？如果可以的话，那么相对于对单一法则进行测试而言，对 250 个法则进行回溯测试不失为一个好的想法。

幸运的是，对于数据挖掘领域而言，对这两个问题的回答是有条件的“是的”。这些条件涉及用于计算法则平均收益率或者绩效统计量的观察资料的数量。如果使用的观察资料的数量足够多，那么对于这两个问题的回答就是肯定的。这对于数据挖掘者来说绝对是个好消息。然而，当观察资料的数量很少时，那么数据挖掘是无法发挥作用的。而我们要想挑选法则就只能靠猜了，在这种情况下，即使对大量的法则进行数据挖掘也不会产生更好的法则。

由于使用的观察资料的数量非常重要，读者无疑想要知道到底需要多少观察资料才够。遗憾的是，这个问题回答起来并不简单。所需要的观察

资料数量取决于原始数据的绩效统计量和本质。我现在所能做的就是提供反映所需样本容量的一般状态的代表性图形。

在前面的图形中，纵轴代表着每年的年化收益率百分比中的数据挖掘偏差。然而，需要注意的是，在图 6-43 中，纵轴代表着被数据挖掘器所挑中的法则的真实优值，或者拥有最高观察绩效的法则的预期收益率。很明显，这种曲线是不能够产生真正的技术分析法则的，因为它们的预期收益率都是未知的。横轴代表着通过测试发现最佳观察绩效的 ATR 的数量。这组数字包括 1（没有数据挖掘）、2、4、8、16、32、64、128 和 256。图中有 3 条曲线。每一条曲线都是以用于计算 ATR 平均年收益率的不同数量的月度观察资料为基础的。用于此案例的月度观察资料的数量设定为 2、100 和 1 000。

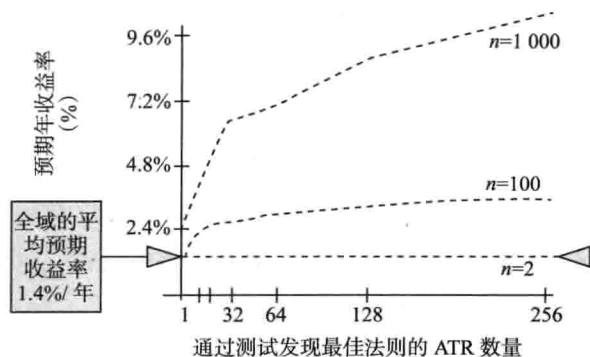


图 6-43 最佳 ATR 绩效的预期收益率 VS. 进行回溯测试的 ATR 的数量

第二个问题是：对大量的法则进行测试就真的能够发现具有更高预期收益率的法则吗？如果真是这样的话，那么预期收益率的曲线就应当随着用于发现最佳法则的 ATR 的数量的增加而上升。换句话说，拥有最佳观察绩效的 ATR 的预期收益率应该与为了发现最佳法则而进行测试的那些法则呈正相关关系。另外，如果曲线没有上升，就意味着过多的研究并不能够导致具有更高优值的法则的出现。这就说明数据挖掘是徒劳的行为。

在图 6-43 中，3 条曲线中的两条出现了上升，即分别以 100 和 1 000 个月份作为观察资料的曲线。这说明进行回溯测试的法则越多，就越容易发现更高优值的法则。以 1 000 为基础的曲线要比以 10 为基础的曲线陡峭

得多。这说明当绩效统计量是从大量的观察资料中计算得出的时候，数据挖掘就越有可能发现更好的法则。这对于关注样本容量的数据挖掘者来说是个好消息。

图 6-43 中没有上升的曲线告诉我们：当用于计算绩效统计量的观察资料太少的时候，数据挖掘是不起作用的。也就是说，对更多的法则进行测试并不能够导致优秀法则的发现。平滑的曲线表明当只用 2 个月的观察资料来计算每个 ATR 的平均收益率时，相对于只对 1 个 ATR 进行测试得出的最佳绩效而言（没有数据挖掘），在 256 个 ATR 中绩效最佳的那个 ATR 的预期收益率也并不高。因此，样本容量的重要性不能被过度强调！

现在让我们回到第 1 个问题中来：数据挖掘所选择的具有更高预期收益率的法则比从全域中随机抽取的法则更好吗？图 6-43 也提到了这个问题。注意，当进行测试的 ATR 的数量增加的时候，以 2 个观察资料为基础的曲线仍旧保持在 1.4% 的水平上。1.4% 的收益率是 ATR 全域的平均 ATR 预期收益率。换句话说，如果我们从全域中随机抽取一个 ATR，并且记录下它的预期收益率，然后再将这一行为重复若干次，那么这个随机抽取的法则的平均预期收益率就是 1.4%。图 6-41 说明当只有 2 个月度观察资料可以使用时，被选中的那个具有最佳绩效的 ATR 还不如随机抽取的那个 ATR 的表现好。曲线是平滑的这一事实说明有多少 ATR 进行测试都是无关紧要的。然而，以 100 和 1 000 个观察资料为基础的曲线的上升程度却超过了全域的平均水平。例如，当使用 1 000 个历史月度资料计算观察到的平均收益率时，从 256 个 ATR 中选出的最佳法则的预期收益率是每年 10%，远远超过了全域的平均值 1.4% 的水平。从本质上讲，当观察资料的数量非常充足时，所观察到的法则的收益率是非常有帮助的，尽管其预期收益率的指示有出现偏差的可能性。然而，当观察资料太少时，观察到的收益率实际上是没有价值的。

将数据挖掘偏差视为全域规模的函数：在可变优值的全域中

前面的内容证实了在充足的样本容量下，随着进行测试的法则数量的增加，数据挖掘也随之增加的结果。这说明最佳法则的观察绩效是很有价

值的, 尽管其预期收益率的标准存在偏差。但问题是: 当全域包含具有不同预期收益率的法则时, 与最佳法则(观察到的最高绩效)有关的偏差到底有多大。

图 6-44 ~ 图 6-47 描绘了数据挖掘偏差(纵轴)与在法则全域中进行回溯测试的法则的数量之间的对比情况, 在法则全域中, 优值是根据图 6-42 进行分布的。每一个图都是以计算平均收益率的不同月度观察资料的数量为基础的, 即 2 个月、100 个月和 1 000 个月。首先, 每一条曲线都单独表示(见图 6-44 ~ 图 6-46), 图 6-47 显示了 3 条曲线叠加在一起的情况。从中我们可以看到在不同的样本容量下, 数据挖掘偏差的相关程度。

图 6-44 仅仅使用 2 个月的观察资料来计算每个 ATR 的收益率。图中的数据清楚地表明由小样本容量导致的问题。例如, 在 256 个 ATR 当中观察到的那个年收益率 200% 的最佳 ATR 的绩效夸大了其预期收益率。对于短期的历史绩效来说, 数据挖掘偏差是个极值。然而, 在图 6-43 中显示的 100 个观察资料的案例, 以及图 6-44 中显示的 1 000 个观察资料的案例中, 数据挖掘偏差的值很小。例如, 当使用 100 个月度资料时, 与 256 个最佳 ATR 有关的数据挖掘偏差是 18%, 而使用 1 000 个月度资料时, 偏差值则小于 3%。

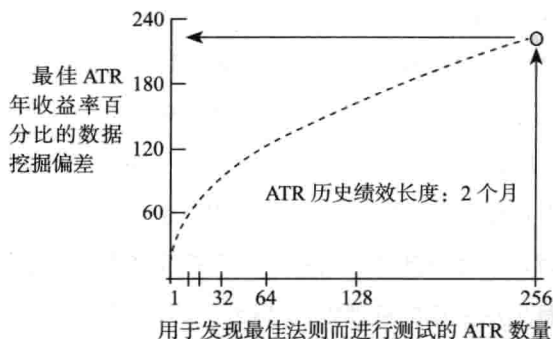


图 6-44 数据挖掘偏差 VS. 用 2 个月的历史绩效测试
在可变绩效法则全域中的 ATR 数量

图 6-44 ~ 图 6-46 传达了两个信息。首先, 观察资料的数量对数据挖掘偏差有着重要的影响。其次, 当用于计算法则绩效的观察资料的数量很大时, 数据挖掘偏差很快就会趋于平稳。如图 6-46 所示, 它是以 1 000 个

月的历史绩效为基础的。图 6-47 把前面的三个数字进行了叠加。它说明一旦进行测试的 ATR 超过适当的数量时，那么用大量的 ATR 来寻找最佳的法则（即没有进行数据挖掘）并没有增加偏差的大小，换句话说，一旦超过法则更低的临界值，那么对于用更多的法则进行测试的惩罚也不会增加了，当然，这是在足够多的观察资料都用于计算绩效统计量的情况下。

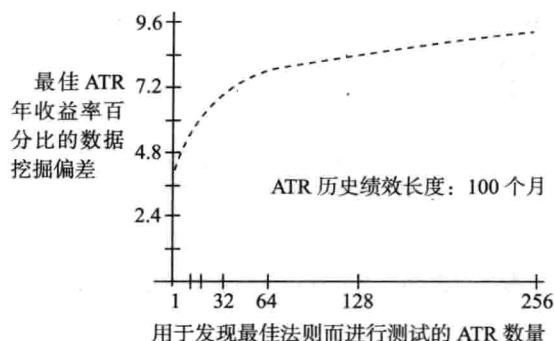


图 6-45 数据挖掘偏差 VS. 用 100 个月的历史绩效测试在可变绩效法则全域中的 ATR 数量

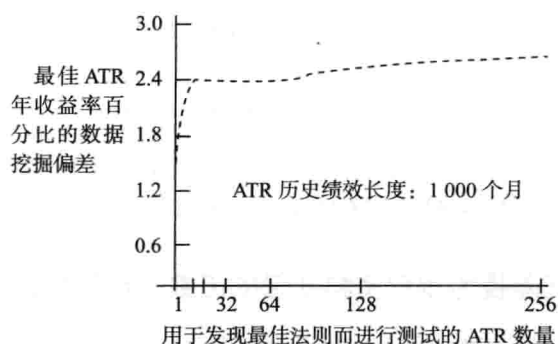


图 6-46 数据挖掘偏差 VS. 用 1000 个月的历史绩效测试在可变绩效法则全域中的 ATR 数量

一旦突破临界值，数据挖掘偏差趋于平稳的速度到底有多快，如图 6-48 所示，这是经过放大的 100 个月的观察资料的曲线图的一部分。它显示了沿着横轴进行测试的 ATR 数量（1 ~ 40）的范围。在对大约 30 个 ATR 进行测试之后，随着进行测试的 ATR 数量的增加，其对数据挖掘偏差的大小的影响是最小的。这对于利用充足的观察资料的数据挖掘者来说是

个好消息。从本质上讲：对大量的法则进行测试是件好事，但是如果用于计算以挑选法则为主的绩效统计量的观察资料十分充足的话，那么这样做所产生的不利后果将会很快减少。

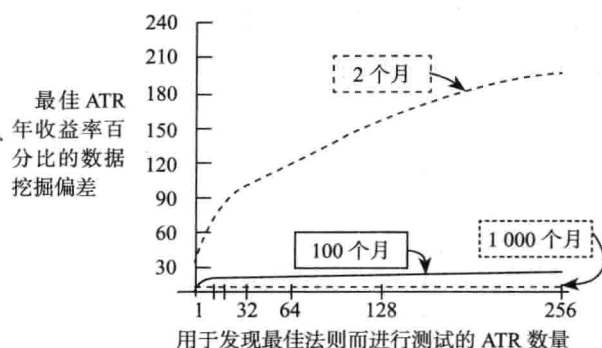


图 6-47 数据挖掘偏差 VS. 在可变绩效法则全域中的 ATR 数量

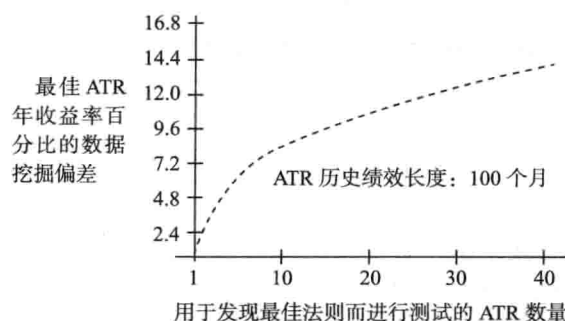


图 6-48 当样本容量足够充足时，数据挖掘偏差将很快趋于平稳

将数据挖掘偏差视为观察资料数量的函数：在可变绩效的全域中

这部分所给出的绘图显示了数据挖掘偏差（纵轴）与用于计算观察到的绩效统计量（横轴）的月度观察资料的数量之间的关系。我们通过之前分别用 2 个月、100 个月以及 1 000 个月的 ATR 数量与数据挖掘偏差曲线的对比来观察其效果。然而，下面这部分图形所描绘的是与若干月份数量的连续统一体相关的观察资料数量和数据挖掘偏差之间的关系。

每一个图形显示了数据挖掘偏差作为观察资料数量的函数的情况，其中的观察资料数量按照从 1 ~ 1 024 个月的顺序排列。我们所期望看到的，在前面的发现的基础之上，看看在观察资料很少的情况下，数据挖

掘偏差是否会扩大。这种情况确实发生了。数据挖掘者从中获得的信息很清晰: 留意大数法则。图 6-49 ~ 图 6-51 显示了数据挖掘偏差与用于计算 3 个不同 ATR 平均收益率 (2 个最佳 ATR、10 个最佳 ATR 以及 100 个最佳 ATR) 的观察资料的数量之间的比较情况。图 6-52 显示了 3 条叠加的曲线。

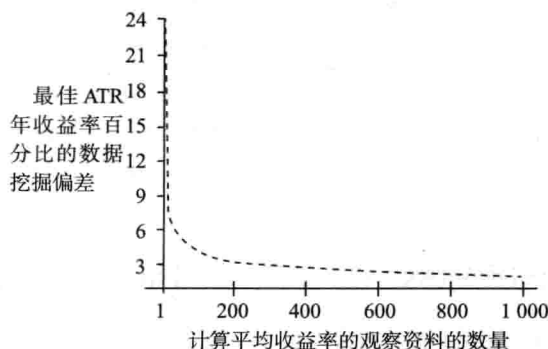


图 6-49 数据挖掘偏差 VS. 用于计算 2 个最佳 ATR 的观察资料的数量

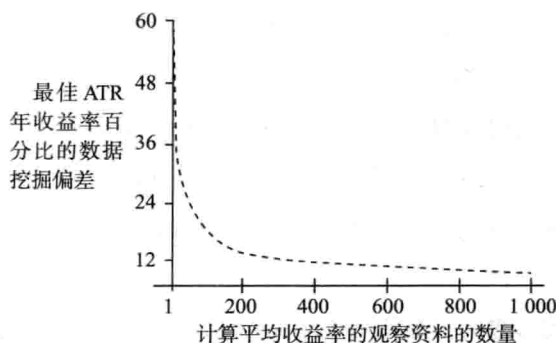


图 6-50 数据挖掘偏差 VS. 用于计算 10 个最佳 ATR 的观察资料的数量

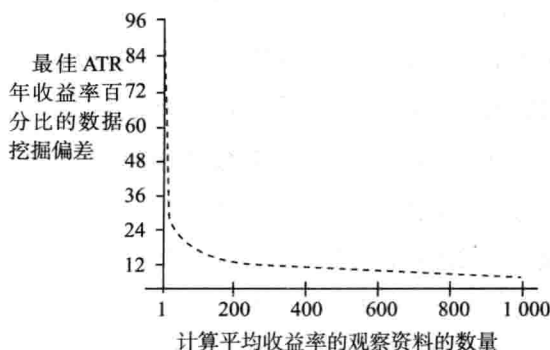


图 6-51 数据挖掘偏差 VS. 用于计算 100 个最佳 ATR 的观察资料的数量

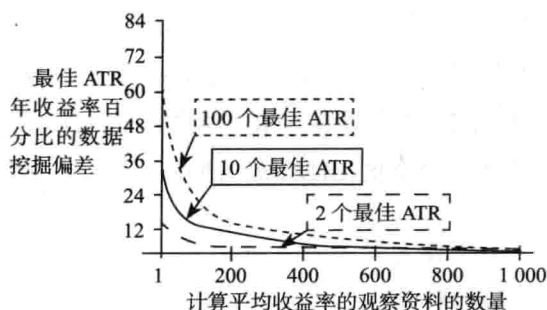


图 6-52 数据挖掘偏差 VS. 用于计算 2 个、10 个以及 100 个最佳 ATR 的观察资料的数量

观察到的绩效与预期绩效 VS. 在可变绩效全域中观察资料的数量

图 6-53 和图 6-54 显示了两条曲线。上面的那条曲线是在 10 个 ATR 当中那个具有最佳绩效的 ATR 的绩效。下面的曲线表示的是那个最佳 ATR 的预期收益率。纵轴代表的是年收益率，既包括观察到的绩效，也包括预期绩效。例如，在图 6-53 中，上面曲线中的 A 点代表着当 400 个观察资料用于计算观察到的 ATR 的平均收益率时，10 个 ATR 中那个最佳 ATR 的绩效。B 点代表着 ATR 的预期收益率。在 10 个最佳法则和 400 个月度观察资料这种特定的条件下，纵轴上 A 点与 B 点之间的距离就是数据挖掘偏差。

我们应该能够想到，这就是之前实验的结果。当观察资料的数量很少时，优值的出现是随机的，且被观察到的绩效将会大于预期收益率，即扩大了的数据挖掘偏差。这一点我们可以从图中两条相隔距离很远的曲线看到。由于单一的 ATR 可以受到幸运的眷顾，所以当绩效统计量是以少量观察资料为基础的时候，两条曲线之间的纵向缺口就会扩大。然而，当更多的观察资料用于计算平均收益率时，且观察到的绩效能够更加真实地代表其预期绩效时，两条曲线就趋向于相互靠拢。但是，由于观察到的绩效中包含有运气的成分，而选择最佳法则的标准又会导致正的偏差，所以，观察到的绩效的曲线通常会在预期绩效的曲线之上。

以上只是对 2 个 ATR 进行测试。图 6-53 显示了对 10 个最佳 ATR 进

行测试的情况。图 6-54 显示了对 500 个最佳 ATR 进行测试的情况。和我们预想的一样，当观察资料的数量增加（曲线靠拢）时，数据挖掘偏差便会缩小。同理，当通过对大量的 ATR 进行检验来发现最佳 ATR 时，由于运气成分的增加，对 500 个 ATR 进行测试得出的偏差比对 10 个 ATR 进行测试得出的偏差要大。

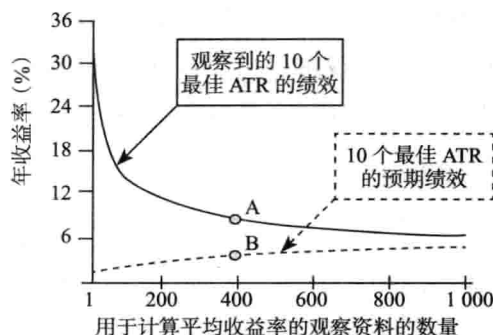


图 6-53 观察到的最佳 ATR 的绩效及其预期绩效 VS. 以 10 个 ATR 为基础的观察资料的数量

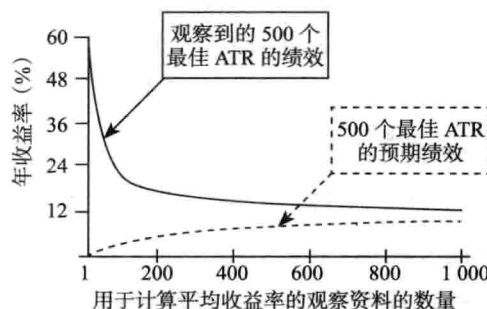


图 6-54 观察到的最佳 ATR 的绩效及其预期绩效 VS. 观察资料的数量

数据挖掘偏差作为法则相关性的函数：以 500 个 ATR 为基础的可变绩效全域

我们之前已经了解到，当备选法则间的收益率相互关联时，就会减少数据挖掘偏差。由于相关性降低了进行测试的法则数量的有效性，所以数据挖掘偏差就会因此而减少。然而在某些极端的案例当中，所有法则的收益率都具有完全的相关性，也就是说，能够用于测试的法则实际上只有 1

个。在这种情况下，就不会发生数据挖掘，而数据挖掘偏差也会因此而降为0。具有同等绩效的法则的全域得出的就是这个结果；它们所有的预期收益率都是0。现在，我们来分析拥有可变绩效全域的情况。

图 6-55 描绘的是当进行分析的 ATR 的数量（横轴）为 1 ~ 256 时，全域中拥有最佳绩效的 ATR 的预期收益率（纵轴）。图中有 4 条曲线，每一条曲线代表着可以从挑选出最佳法则的 ATR 的收益率的不同相关性水平。我们把 4 个相关性水平设定为 0、0.3、0.6 和 0.9，并且把用于计算 ATR 平均收益率的月度观察资料的数量设定为 100 个月。

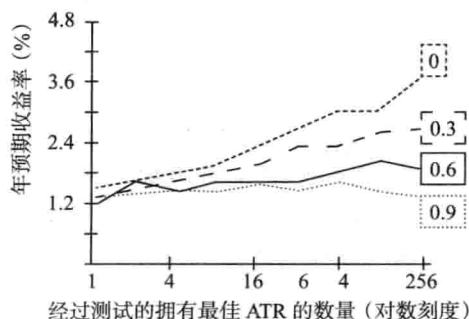


图 6-55 最佳 ATR 的预期收益率 VS. ATR 的数量。在 4 个水平上进行测试：0、0.3、0.6 和 0.9。观察资料数量=100

从图 6-55 中我们看到，相关性水平越高，拥有最佳绩效的法则的预期收益率就越小。这与我们之前得到的进行回溯测试的 ATR 越少，发现最佳法则的概率就越小的这一结果相一致（见图 6-43）。高度的相关性意味着通过更小的全域要比经过测试的实际 ATR 的数量所显示的信息更有效地进行最佳法则的寻找。正如我们之前所指出的那样，在那个极端的案例中，相关性为 1.0，因此也就不存在数据挖掘了。当 ATR 的收益率之间的相关性接近理想的数值 0.9 时，最佳的 256 个 ATR 的预期收益率仅为 +1.4%。这与随机提取的 ATR 的预期收益率是相同的。换句话说，在相关性为 0.9 的水平上，经过测试的有效的 ATR 数量非常小（少量的数据挖掘）。数据挖掘进行的越少，意味着发现优秀法则的概率越小。相反，当收益率之间的相关性等于 0 时，有效的全域规模达到最大化——我们确实在 256 个不同的 ATR 当中进行寻找。因此，256 个 ATR 当中的那个最佳 ATR 的预期收

益率接近 4%，远远超过全域的平均值 1.4%（数据挖掘起作用了！）。

关于数据挖掘偏差研究结果的总结

这部分内容介绍了大量实验的结果。我们现在对研究结果进行总结是非常有价值的。这些结果对那些必须通过某些方法才能解决的问题的本质做出了解释，这些方法我们将在接下来的内容里讨论。

（1）由数据挖掘器发现的法则的观察绩效，即在由进行回溯测试的法则组成的集合中产生的最佳绩效，肯定会产生数据挖掘偏差。其预期收益率小于其已经观察到的历史绩效。而且其在未来的任何数据样本中可能出现的绩效都要低于通过回溯测试产生的可以赢得绩效竞争的绩效。

（2）用于计算绩效统计量（用来挑选最佳法则）的观察资料的数量对数据挖掘偏差的程度有着深远的影响。观察资料的数量同样影响着绩效统计量的抽样分布的宽度和尾厚度。

a. 观察资料的数量越多，抽样分布就越窄，数据挖掘偏差也就越小。

b. 观察资料的数量越多，抽样分布的尾部就越轻，数据挖掘偏差就越小。

（3）寻找到的法则的数量越多，数据挖掘偏差就越大。

（4）存在于被测试的法则中真实品质（预期收益率）的变异性越小，数据挖掘偏差就越大。换句话说，进行回溯测试的各法则的预测力越相等，数据挖掘偏差就越大。但是这一点并没有经过实验的考察。

（5）当样本容量非常大的时候，数据挖掘才能发挥效力。进行测试的法则越多，最终被挑选出来的法则预期收益率就越高。而当样本容量非常小的时候，数据挖掘则不发挥作用。

解决方法：处理数据挖掘偏差

数据挖掘虽然发挥作用，但是肯定会产生偏差。因此，如果通过数据挖掘器进行分析的备选法则的集合包括任何具有更优预期收益率的法则的话，那么就可以使用数据挖掘来发现这些法则。然而，它们的预期收益率

很有可能低于它们在回溯测试中所产生的绩效。样本外绩效的逐渐降低已经成为了数据挖掘工作中无法改变的事实。

当所有的备选法则不再具有价值的时候，更加严重的问题就出现了，即这些法则的预期收益率就会等于或者小于0。在这种情况下，数据挖掘器就会具有受到数据挖掘偏差误导的危险，并会挑选出一个预期收益率为0的法则。这就是客观的技术分析师挖到的金矿。我们这部分的内容将要讨论降低这种风险的方法。

目前一共提出了3种方法：样本外测试、由马科维茨（Markowitz）和 Xu 提出的数据挖掘校正因子以及随机化方法。样本外测试需要把一个或者多个历史数据子集从数据挖掘（样本外）中排除出去。这个数据将用来对在数据进行挖掘中（样本内）发现的最佳法则进行评估。我们把在当前数据中所选法则的绩效与在未来的数据中提供该法则预期收益率无偏估计的挖掘运算相隔离。我们将历史数据分割成若干的样本内和样本外片段。我在后面的内容中对此进行了总结。

第二种方法是一种以诸如自举法和蒙特卡罗等方法为基础的随机化方法，但是这种方法还没有在学术文献中进行广泛的讨论。该方法具有以下几点优势。首先，它使得数据挖掘器可以对任何数量的法则进行测试。对大量的法则进行测试，会加大发现更优法则的概率。其次，它不需要像在样本外测试那样把数据放在一边，因此，也就允许我们把所有能够获得的历史数据全部用于数据挖掘。最后，它使我们可以进行显著性检验，有时还可以产生大致的置信区间。

第三种方法是由马科维茨和 Xu 提出的数据挖掘校正因子，它降低了那个表现最佳的法则的观察绩效。读者可以把精力集中在这种方法的原始资料上面。关键在于，用这种方法进行的有限的实验表明它不仅可以取得让人出乎意料的效果，也会让人感到无比沮丧，当然，这完全要依靠当时的情况来决定了。

样本外测试

样本外测试是以在样本外数据中进行数据挖掘的法则的绩效可以提供

其未来绩效的无偏估计这一有效理论为基础的。³³ 当然, 由于抽样变异的原因, 法则在未来的表现可能与其在样本外测试中的表现有所不同, 但是, 我们也没有理由因此就认为其表现就会更差。因为样本外绩效是无偏的, 所以为样本外测试保留数据也是合情合理的。

由此我们会面对如何将历史数据分割为样本内子集和样本外子集这一问题。我们提出了各种不同的办法。最简单的方法就是创建两个子集: 用于数据挖掘的历史数据的早期部分和为样本外测试保留数据的后期部分(见图 6-56)。³⁴ 更多复杂的分区在棋盘状的图形中将历史数据进行拆分, 这样, 所有的样本内和样本外数据就会来自所有的历史部分, 如图 6-56 所示的图形 B 和 C。

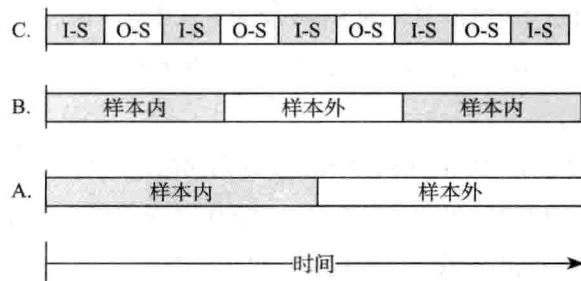


图 6-56 历史数据的样本内和样本外分割

1. 前推检验

另外一种针对金融市场交易应用的数据分割方法是前推检验法。这种方法是由帕尔多 (Pardo)、³⁵ 德·拉·马萨 (De La Maza)、³⁶ 卡茨 (Katz) 和麦科密克 (McCormick)³⁷ 以及考夫曼 (Kaufman)³⁸ 联合提出的。该方法使用了将自身分为样本内部分和样本外部分的移动数据窗口。然而, 用于样本内和样本外的分段的术语在某种程度上是不同的。因为前推检验是动态的, 所以法则会随着市场的不断变化而得到修正, 这个术语也暗中说明了法则是从这些经验中得出的。因为法则的参数值是通过这部分数据获得的, 所以样本内的分割是指序列 (training) 的数据集合。由于从序列测试中获得的参数值要在这个数据分割中进行绩效的测试, 所以样本外的分割被称为测试 (testing) 数据集合。

在时间序列分析和技术分析中,移动数据窗口的概念已经广为人知了。移动平均数、移动最大-最小价格通道以及变化率指标等,都提供了一个随着时间轴滑动的数据窗口。比如决定最大值或者平均值的数学运算就属于数据窗口的范围。

在这种情况下,进行前推的数据窗口(序列集合+测试集合)有时是指重叠的意思(见图6-57)。我们计算每一个重叠部分的样本外绩效。由于存在大量的重叠,因此就要计算大量的样本外绩效,这使我们有可能确定绩效统计量的方差。同时,置信区间也是可以进行计算的。此外,用移动数据窗口和动态自适应对法则进行测试。那么在理解了每一个新数据窗口之后,一个新的法则就会形成。前推检验对于金融市场等不稳定的现象来说是具有特殊吸引力的。Hsu 和 Kuan 发现了自适应交易法则(他们把这些称为学习策略),这在他们对超过 36 000 个客观的法则进行测试的过程中非常有效。³⁹

在前推检验结构中,序列分割用于数据挖掘,然后把从中发现的最佳法则用于测试分割中。由此便产生了该法则样本外绩效的无偏估计。然后,再将整个的窗口向前推移,并将这一过程进行重复。

前推检验的过程如图6-57所示。需要注意的是,向前推移的窗口要足够的远,不能出现独立地进行测试的数据分割不能重叠的情况。通过这种方法,计算出来的样本外绩效才能相互保持独立。

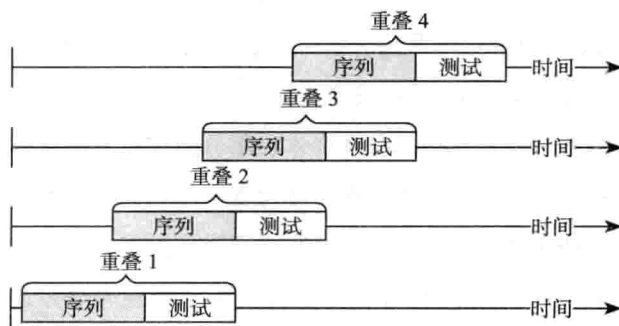


图 6-57 运用前推检验数据对两个数据分割进行测试

2. 数据样本外测试方法的局限性

当对数据挖掘偏差进行校正的时候,样本外测试受到了很多缺陷的困

扰。第一也是最重要的一点是，为样本外测试保留的数据具有寿命短的特征。只要这些数据被使用过一次，它们就失去了价值。所以，这些数据也就不可能再提供对法则绩效的无偏估计了。

样本外测试的第二个缺陷是它消除了在数据挖掘中某些数据的比例。因此，它减少了能够用于发现形态的数据总量。当噪声很大并且信息含量很少的时候，更多得到的观察结果就会是极端的数值。

第三，关于如何分配样本内集合与样本外集合之间的数据比例的决定是随意的。目前还没有一种理论对究竟应该给序列集合和测试集合分配多少数据做出说明。我们对这些选择的结果是非常敏感的。通常来说就是凭直觉来做这些事。

马科维茨 /Xu 提出的校正因子

一种由马科维茨和 Xu⁴⁰ (MX) 提出的评估数据挖掘偏差的公式，对由德·拉·马萨⁴¹制作的电子数据表格的实际实施情况做出了描述。马科维茨在现代证券投资组合理论所做的创举是众人皆知的，他也因此而荣获了1990年的诺贝尔经济学奖。

MX 理论并不需要数据段。它所需要的是在挑选最佳绩效法则的过程中，通过数据挖掘器进行分析的所有法则的完整区间收益率（每日、每周、每月等）的集合。该方法通过面对全部经过测试的法则的平均绩效时，收缩观察到的最佳法则的绩效的方式来校正数据挖掘偏差。收缩的程度主要由以下几个因素构成：①围绕在全部经过测试的法则的每日平均收益率周围的每一个法则的每日收益率的变动程度；②围绕在所有经过测试的法则的平均收益率周围的每一个法则的平均收益率的变动程度；③进行分析的法则的数量；④收益率的时间区间（例如，天数）的数量。

在德·拉·马萨早期发表的论文中，就已经开始设计用于计算收缩因子 B 的电子表格的具体细节了，“ B ”是指 在 0 和 1 之间的数字。我们把这个数值代入下面的公式：

$$H' = R + B(H - R)$$

其中， H' 表示在对数据挖掘偏差进行调整之后的具有最佳绩效的法则

的预期收益率； R 代表所有经过测试的法则的平均收益率； H 表示观察到的最佳法则的绩效； B 代表收缩因子。

对最佳法则进行的经过校正的与缩小了的绩效评估 H' ，将会位于其已经观察到的绩效与全部经过测试的法则平均绩效之间的位置。当 B 等于 0 时，最佳法则的预期收益率仅仅等于所有法则 (R) 的平均收益率。当 B 达到数值 1 时，最佳法则的预期收益率等于其在回溯测试中得到的收益率。这种情况只有在没有检测到数据挖掘偏差的情况下才会出现。因此，也就不需要对观察到的最佳法则收益率进行收缩了。

由 MX 给出的结果是非常模糊的。在有些情况下，经过缩小的绩效评估是非常有好处的。然而，在实践中却很容易碰到其他的一些情况，也就是说其结果会同现严重的错误。 MX 是模糊原则最好的使用者。

随机化方法

在第 5 章中，我们介绍了两种用于在回溯测试中对观察到的法则的绩效进行统计显著性的方法，即自举法和蒙特卡罗置换法。正如我们在第 5 章中所指出的那样，这些显著性测试中有的法则在对单一的法则进行回溯测试时才会有效，也就是说，我们不需要通过数据挖掘来发现最佳法则。然而，在描述这部分修正的内容时，这些过程也可以用来对通过数据挖掘发现的法则的统计显著性进行测试。

两种随机化方法都允许数据挖掘器使用所有的历史数据来发现法则。这样一来我们就可以避免对样本外测试所需要的数据数量进行随意决定的问题了：如何分配样本内和样本外数据的比例，如何分配用作序列和测试比较的数据数量等。然而，这两种随机化的方法需要保留在回溯测试过程中经过测试的每一个法则有关的信息。因为根据使用自举法和蒙特卡罗置换法的不同，每种方法都必须保留特定的信息，至于每种方法所需要的信息，我们将分别进行讨论。

为了解释随机化方法是如何进行显著性测试的，我们先来进行一个简短的回顾。检验统计量的抽样分布是以统计推断为基础的。它描述的是在统计量当中随机化的程度。这个抽样分布对于判断被观察到的法则绩效的

统计显著性是非常重要的。第5章还指出，抽样分布可以从自举法和蒙特卡罗置换法当中衍生出来。

1. 自举法抽样分布：怀特的真实检验

怀特的真实检验（White's Reality Check, WRC）由位于圣迭戈的加利福尼亚大学的经济学家赫伯特·怀特（Halbert White）教授发明并获得专利保护的。他使用适用于测试通过数据挖掘发现的最佳法则的统计显著性的自举法⁴²推导抽样分布。这款软件可以通过由怀特博士所有的Quantmetrics⁴³公司获得。

在怀特的真实检验以前，自举法曾经应用于用产生的抽样分布测试单一法则的显著性。怀特的发明使他获得了专利权，并让自举法应用到对由数据挖掘发现的法则的测试当中。具体来说，WRC使得数据挖掘器可以对最佳的 N 个法则进行抽样分布，其中的 N 代表在所有法则的预期收益率为0的情况下，进行测试的法则的数量。换句话说，WRC产生的抽样分布可以对原假设（指在数据挖掘过程中所有进行分析的法则的预期收益率等于0）进行测试。

我们用更加简化的例子来解释运用WRC的过程。假设数据挖掘器只对期限为10天的2个法则（ $N=2$ ）进行回溯测试。我把这2个法则设定为 R_1 和 R_2 。在这10天中， R_1 的平均每日收益率为1%，而 R_2 的平均每日收益率为2%。数据挖掘器将 R_2 作为最佳法则挑选出来，但是随即对此产生了疑问，如果2%的收益率是由数据挖掘偏差所导致的，且 R_2 的真实预期收益率实际为0，那怎么办呢？为了产生适合于此次数据挖掘风险的抽样分布，我们按照下面的步骤来进行。

（1）法则测试的是10个交易日的数据。我们用10张纸片代表10个交易日，并按照从1到10的顺序将它们放在抽样的容器中。

（2）将这些纸片采样并更换10次。每一次将一个日期从中取出，并记录下日期的编号。因为这是每一个法则的历史绩效的长度，所以抽样的数量就是10个。这是自举定理所需要的必要条件。假设被抽到的日期排列如下：5、3、5、5、7、2、5、8、1、2。

（3）我们使用步骤2中获得的日期，以这些日期的实际每日收益率为

基础组成 R_1 。我们用同样的方法组成 R_2 。现在，每一个法则就拥有了一个包含 10 个随机产生的绩效历史。

(4) 回想一下我们的目的就是产生在预期收益率都为 0 的法则的全域中出现的那个具有最佳绩效的法则的抽样分布。因为这是可以实现的，实际上也是我们的希望之所在，即至少有一个经过测试的法则具有大于 0 的预期收益率， R_1 和 R_2 的原始收益率必须调整为平均收益率为 0。为了完成这个目的，WRC 决定着每一个经过测试的法则的平均每日收益率，并且把这个总量从每个独立的每日收益率⁴⁴中减去。在这个案例中， R_1 的平均每日收益率是 1%，而 R_2 则是 2%。因此，1% 就会从 R_1 中的每一个每日收益率中减去， R_2 的步骤同 R_1 。这就是改变每个法则的每日收益率的分布，并使其集中度大于 0 的效果。这些经过修正的每日收益率正在被我们所使用。

(5) 经过调整的 R_1 (日期的排列为 5、3、5、5、7、2、5、8、1、2) 的收益率是平均值。 R_2 的步骤同 R_1 。

(6) 两个平均收益率中最大的那个值将会被保留下来成为用于组成 N 个法则之最大平均收益率的抽样分布的第一个数值。在本案例中， $N=2$ 。

(7) 重复步骤 2 ~ 步骤 6 若干次 (比如 500 次或更多)。

(8) 统计量的抽样分布，即在预期收益率为 0 的 $N(2)$ 个法则的全域中的最大平均收益率，是根据这 500 个数值组成的。

(9) 我们利用现在手中的抽样分布来计算概率值。

(10) 概率值就是从步骤 7 当中获得的超出被测试的法则平均收益率的那 500 个值的部分。(过程请参见第 5 章。)

我们将上面所描述的案例有意识地进行高度简化。在一个更加现实的案例中，数据挖掘器将会对时间跨度更大、数量更多的法则进行分析。但是产生抽样分布的过程都是一样的。一旦软件提供了所有进行测试的法则的区间收益率 (每天)，所有这些步骤将全部由 WRC 软件自动完成。这个数据集包括所有 WRC 需要用于推导出适合特定数据挖掘风险的抽样分布的信息。当很多的法则用于测试时，我们就需要大量的数据。目前，只有少数的数据挖掘器和数据挖掘系统能够存储这些有价值的信息。

在本书的第二部分中, 我将对用于标准普尔 500 指数交易的 6 402 个法则进行案例研究。我们在后面的内容里还要讨论一种加强版本的 WRC, 并可以把它作为测试观察到的法则绩效的统计显著性的一种方法。

2. 蒙特卡罗置换法

蒙特卡罗置换法 (MC) 同样可以用来产生需要评估由数据挖掘器发现的法则的统计显著性的抽样分布。正如我们在第 5 章中所说, 它使用不同于自举法的数据测试以及不同的原假设构想, H_0 。

根据数据, MC 运用某法则的输出值 (+1, -1) 的历史时间序列和市场的原始收益率。相反, 自举法使用的是法则的区间 (每日) 收益率。根据原假设 (H_0), WRC 和 MC 都对所有的法则在回溯测试期间不具有预测力的概念进行了测试。然而, MC 对原假设 (H_0) 的阐述有所不同。为了对缺少预测力的法则的绩效进行模拟, MC 用每日市场价格的变化随机配对了对法则的输出值。反观 WRC, MC 对输出值和市场变化的配对没有出现替代的现象。根据推测, 如果法则的输出值与市场变化是随机分配的, 那么其绩效就与认为法则没有价值的原假设相一致。通过随机配对获得的收益率成为了反驳与其相比较的实际法则收益率的基准。如果法则本身确实具有预测力, 那么其确有依据的输出值与市场变化的配对就会产生显著的更优绩效。

随机配对产生的绩效就是用法则的值 (+1, -1) 乘以当天随机配对的 市场变化。因此, 如果法则的输出值 +1 与当天上涨的市场相配对的话, 那么法则获得的就是市场价格变化带来的收益。同理, 如果输出值为 -1 与当天下跌的市场配对, 其收益率就会为正值。当法则的输出值的信号与市场价格变动的信号相反的时候, 亏损就会出现。

当法则的全部历史输出值与市场价格变化配对完成以后, 我们进行平均收益率的计算。由 N 个法则组成的全域中每一个法则都要计算其平均收益率。我们把其中最高的平均收益率挑选出来, 并使其成为第一个用于计算 MC 抽样分布的数值。我们将这套完整的程序重复若干次 (大于 500 次), 每次都用随机选择的不同数值进行配对。这样, 我们就可以通过大量的数值来产生抽样分布了。

需要注意的是,当 WRC 进行这些程序时,我们就不需要对集中于超过数值 0 的每一个法则的收益分布进行重新定位了。相反,MC 的抽样分布的平均值就是在数据挖掘风险中没有价值的法则的预期收益率。如果在将法则的值与市场收益率进行随机配对之后,抽样分布就会自然而然地位于大于 0 的数值上。概率值可以直接从产生 MC 的抽样分布中计算得出,正如我们使用替代方法得出的抽样分布那样。概率值的位置位于分布右侧尾部的区域,而这一区域又位于或者超过经过回溯测试的法则的平均收益率。

相比于 WRC,MC 方法的局限性在于它不能够产生置信区间,这是因为它没有对有关法则平均收益率的假设进行测试。MC 的原假设很简单,就是所有经过测试的法则的输出值与未来的市场行为只是偶然相关。

现在来回顾一下,蒙特卡罗置换法(MC)通过数据挖掘器所产生的抽样分布的步骤如下。

(1) 获得在回溯测试风险中进行测试的所有 N 个法则的每日输出状态。

(2) 每一个法则的输出值与无序的实际市场价格变化进行随机配对。注意同样的配对适用于所有法则的情况。因此,如果法则 7 日的输出值与 15 日的市场收益率相配对的话,那么同样的配对就必须应用于所有的竞争规则。这样做的目的是保护可能会在法则中出现的相关结构,这点是影响数据挖掘偏差的 5 个因素之一。

(3) 确定每一个法则的平均收益率。

(4) 在 N 个法则以外,选择一个最高的平均收益率。这个值将成为在 N 个无效法则之中那个最好的平均值的抽样分布中的第 1 个值。

(5) 将步骤 2~步骤 4 重复 M 次($M=500$ 或者其他大数)。

(6) 根据从步骤 2~步骤 5 获得的 M 值组成抽样分布。

(7) 关于经过回溯测试的法则的概率值可以通过从步骤 6 中获得的值来确定,这个值等于或者大于其平均收益率。

3. 初期 WRC 与 MCP 方法的潜在缺点

这两种方法都可以对认为在数据挖掘全域中的法则都是无效的原假设

进行测试。数据挖掘器希望通过支持备择假设的方式反驳下面的假设: 并不是所有经过检验的法则都是无效的。正如我们之前所指出的那样, 术语无效的意义是取决于所用的随机化方法的。在 WRC 中, 无效法则的预期收益率等于 0。而在 MC 中, 无效是指法则的输出值与市场的未来变化之间的随机配对。

在第 5 章中我们曾经讨论过, 有两种途径可以导致假设检验出现错误。类型 I 的错误是在原假设是正确的(所有经过测试的法则都是无效的)而被错误所排斥的情况下发生的。当一个实际上无效的法则由于运气的原因, 试图去获得高平均收益率和低概率值的时候, 这种错误便发生了。这就是我们所说的数据挖掘者的金矿。

类型 II 的错误是在原假设实际上是错误的, 并且应当被排斥, 但是法则的绩效却太低而不能将其排斥时发生的。也就是说, 该法则实际上具有预测力, 但是其价值还没有被发现, 这是因为法则在回溯测试中的运气不佳。数据挖掘器以错过真正的技术分析金矿而告终。

假设检验的作用就是避免把类型 II 的错误视为测试的能力, 不要把其与法则的预测力相混淆。在此, 我们谈到用统计检验的力量来发现错误的原假设 (H_0)。检验测试具有两个优度度量的特点: 显著性(造成错误类型 I 的概率)和其能力(避免类型 II 错误的概率)。至于初期的 WRC 和 MC 版本的问题欢迎大家提出批评。

经济学家彼得·汉森 (Peter Hansen)⁴⁵ 指出, WRC 甚至是 MCP, 在数据挖掘全域中包括比基准还要差的法则时, 都会出现能力失灵的现象。这实际上说的是一个或者多个比无效法则还要糟糕的情况, 即它们的预期收益率要小于 0。例如, 如果我们想要保留一个预期收益率真正大于 0 的法则的输出值的话, 这种情况就会发生。由于在本书第二部分中进行测试的法则中, 有大约一半的法则是通过这种方式获得的(例如, 反转法则), 所以在第二部分中介绍的研究对于汉森讨论的问题来说是容易受到批评的。

为了对汉森的批评进行评价, 蒂莫斯·马斯特斯 (Timothy Masters) 博士在包含两个优于或者低于基准的 ATR 全域中, 对 WRC 和 MCP 两种方法的能力做出了评估。换句话说, 评估对这个困扰汉森的问题进行了精确

的模拟。这些测试表明，他对具有负预期收益率的法则可以获得巨大差异的收益率的担忧是合理的。这个收益率可以彻底降低 WRC 和 MCP 两种方法的能力。然而，如果在对技术分析法则进行数据挖掘的过程中遇到这种情况，两种方法就可以合理地抵抗这个问题。因此，当在进行分析的法则的全域中出现一个或者更多的更优法则时，WRC 和 MCP 两种方法都有发现它们的机会。

4. 近来对 WRC 和 MCP 的改善

不管 WRC 和 MCP 两种方法的初始版本的合理性如何，最近由罗马诺 (Romano) 和沃尔夫 (Wolf)⁴⁶ 发表的论文介绍了对增强 WRC 与 MCP 能力所进行的修正。因此，罗马诺和沃尔夫所做的强化工作降低了出现类型 II 错误的概率。这里提到的修正将被介绍到用于第二部分的案例研究中的 WRC 与 MCP 版本中。然而，现在没有一款经过强化的 WRC 的商业版本。而加强版的 MCP 已经可以在公共领域看到，该版本的密码可以在 www.evidencebasedta.com 网站中获得。

第7章 Chapter 7

非随机价格波动理论

如果市场的波动是完全随机的，那么技术分析就是毫无意义的，其在表述上所面临的风险也是显而易见的。当且仅当价格的运动在某种程度上以及在某时间段内是非随机的话，则技术分析就是一种合理的尝试。

尽管对非随机价格波动的出现进行证明是必要的，但是这种单独的证明并不足以证明任何具体的技术分析方法。每一种方法必须通过客观实证来证明它捕捉某阶段非随机价格波动的能力。我们将在第二部分对大量的技术分析法则在这些方面的表现进行评估。

理论的重要性

很多来源于行为金融学领域的新兴理论对“为什么价格运动在某种程度上是非随机的，以及其为什么具有潜在预测力”进行了解释。因此，这些新理论所采取的立场与 40 多年来一直被视为金融理论奠基石的有效市场假说（EMH）是对立的。有效市场假说主张金融市场中出现的价格变动都是随机的，是不可预测的。所以，技术分析对这些新兴的理论是抱有很大希望的。

读者也许会感到奇怪，为什么对非随机性的理论支持会如此重要呢？如果一种技术分析方法在回溯测试中能够获得盈利，那么我们还需要理论支持吗？可以肯定的是，在回溯测试中产生的显著盈利不仅可以证实市场是非随机的，而且还证实了该法则可以利用市场非随机性的某个部分。

这种观点忽视了理论支持的重要性。即使采取所有的统计预防措施，但是那个不能被强力支持的具有盈利性的回溯测试仍然是一个孤立的结果，

并且被视为含有运气成分。成功的回溯测试通常会引起这样一个问题：该法则会在未来依旧表现上佳吗？理论之所以可以有力地帮助我们，是因为与强力理论相一致的回溯测试是不可能再次拥有统计学上的巧合的。当回溯测试具有理论基础时，它就不再是孤立的事实，而是一幅更大的相互融合的图片中的一部分，其中，理论对事实进行解释，而事实又反过来证实理论。

例如，在商品期货中的技术趋势跟踪系统的盈利性可以被解释为对提供给商业套期保值者有价值服务的补偿（风险最小化），也就是风险转移。换句话说，经济理论预测期货市场应当足以证明具有盈利性的趋势，来引导趋势追踪者接受价格风险套期保值者所希望下跌的程度。¹

科学的理论

在日常谈论中，常有关于“为什么事情本来就应该是这样”的说法。就像我们在第3章中解释的那样，科学理论就是与众不同的东西。首先，它对以前大范围的观察资料做出了简要的解释。其次，也是最重要的一点，科学理论做出了在以后被新观察资料所确认的具体预测。

开普勒的行星运动定律（laws of planetary motion）² 简要地解释了之前大量的天文学观察结果，并且对后来由新增加的观察资料所证实的事物做出具体的预测。然而，和开普勒定律一样，牛顿的“万有引力原理”取得了更好的成绩，因为他观察的范围比开普勒更大。牛顿的“万有引力原理”不仅解释了开普勒定律的工作原理，而且还对更多现象做出了解释。非随机价格波动理论不需要用精确的物理理论做出预测。然而，为了使其更加有效，非随机价格波动理论不仅对以前观察的众多市场行为进行简要的描述，而且还要对后来的观察资料所证实的事物做出可以检验的预测。

流行的技术分析有什么问题

很多关于技术分析的教科书并没有对书中提到的方法是如何工作的做

出解释。技术分析开创性作品³的作者之一，约翰·迈吉的陈述就很具有代表性。他说：“我们不希望知道市场为什么会如此运行，我们只希望理解它是如何运行的就够了。历史具有很明显的趋势重复性，知道这些就足够了。”

有些技术分析的教科书确实给出了解释，但是这些具有代表性的特殊理论只能产生不可检验的预测。根据约翰·墨菲（John Murphy）的观点，技术分析的基本前提是，“任何可以影响价格的事情（根本的、政治的、心理的，或者其他方面）确实通过市场的价格所体现出来的因素”。⁴我们很难从这些模糊的表述中提取出可以检测的预测。

实际上，这些声称可以解释技术分析工作原理的表述存在逻辑上的矛盾。例如，如果价格确实反映（打折）所有的可得信息，那么这就暗示了价格是缺少任何预测信息的。为了理解其中的缘由，试想一个假设的价格形态出现在股票 XYZ 中，XYZ 当前的股价是 50 美元，而该图形形态则暗示着其价格会达到 60 美元。然后，根据定义，当这个形态出现时，其具有预测性质的暗示则还没有反映到股价当中来。然而，如果价格反映所有信息这一前提是真实的，那么价格就应当已经处于预测的水平上，因此，就会降低该形态的预测能力。

当我们提到有效市场假说的时候，这种逻辑上的矛盾就会更加明显，有效市场假说是一种拒绝技术分析的有效性，它是建立在“价格完全反映所有可得信息”⁵这一前提之上的。技术分析和它的对立者是不可能以同样的前提条件作为基础的。这种矛盾就是非批判性思维的样本，这在当下流行的技术分析版本中是常见的。

幸运的是，技术分析的“价格完全反映所有可得信息”这一基础前提似乎与事实是矛盾的。其中一个实例就是由行为金融学提出的所谓“滞后反应”的假设。由于价格有时不能对新的信息做出像有效市场假说理论所主张的那样快速的反应，即当新信息出现时，确实反映其价格水平的系统价格运动，或者趋势。价格未能快速对信息做出反应是由于一系列困扰着投资者的认知性错误所导致的，比如保守主义偏见和锚定效应。我们将在后面的章节对这些内容进行讨论。

另外一种对于当下流行的技术分析所做的辩护是以流行的心理学为基础的。说到流行的心理学，我的意思是指那些广受欢迎，但是又缺少科学支持的人类行为准则，根据在行为金融学领域的著名经济学家和权威——罗伯特·希勒（Robert Shiller）所说的，在从心理学中得到的大量教训中，要注意很多流行的心理学投资根本就是不可靠的。在市场繁荣期间，投资者被认为是愉快的和极度激动的；而在市场崩盘期间，投资者又表现出惊慌失措的情绪。在这两种市场环境下，投资者就像盲目地追随着羊群的绵羊一样，根本不会有他们自己的想法。⁶事实上，人们要比这些流行的心理学理论所表达的观点理性得多。在那些最有影响力的金融事件中，大多数人都专注于其他人的个人事情，而不是关注金融市场本身，所以，我们很难想象市场可以完全反映由这些心理学理论所描述的情绪。⁷为了证明技术分析，我们需要关注那些超越流行的技术分析教科书以外的东西。幸运的是，在行为金融学以及其他领域得到发展的理论正在提供技术分析所需要的理论支持。

反对者的观点：有效市场和随机游走

在对“为什么非随机价格运动应该存在”这一理论做出解释之前，我们需要分析一下反对者的观点，即有效市场假说。近来，有人在辩论有效市场假说不一定能够预示价格会跟随不可预测的随机游走的趋势，⁸而且有效市场和价格的可预测性是可以共存的。然而，有效市场假说的前辈主张随机游走是有效市场的必然结果。这部分内容会对这些情况进行陈述，并且分析其缺点。

什么是有效市场

有效市场指的是不能够被打败的市场。在这个市场中，没有任何基本的分析和技术分析策略、公式或者系统可以获得经过风险修正之后还能够跑赢大市（由基准指数所定义的）的收益率。如果市场确实是有效的，那么通过买入和卖出市场指数获得的经过风险修正的收益率

就是我们希望看到的结果。这是因为，有效市场中的价格可以完全反映所有已知和可知的信息，所以，当前的价格为每一个证券的价格提供了最佳评估。

根据有效市场假说，市场之所以会达到一个有效的价格状态，是因为很多理性的投资者试图为了将其财富最大化而付出的不懈努力的结果。在追求真实价值的过程中，这些投资者以一种相对正确的方式⁹利用最新的信息来更新他们自己的观点，并对证券未来的现金流做出预测。尽管没有哪个单独的投资者是无所不知的，但是很多的投资者聚在一起就可能知道尽可能多的知识。这些知识会促使投资者在均衡或者理性的价格水平买入并卖出证券。

在这样的世界里，只有在获得新的信息时，价格才会出现变化，而且价格几乎会马上达到能够完全反映信息的新的理性价格水平。因此，价格走势并不是按照从一个理性价格向下一个理性价格逐步发展的，这一点也得到了趋势分析师的认同。根据有效市场假说理论，当价格快速地从一個理性价格到达下一个理性价格时，价格的轨迹实际上就是一个阶梯函数。当获得利好消息时，股价就会上升；反之，则会下跌。由于无法预测利好与利空消息（如果是已知的，就不能称为消息了），因此价格的变化也无法预测（见图 7-1）。

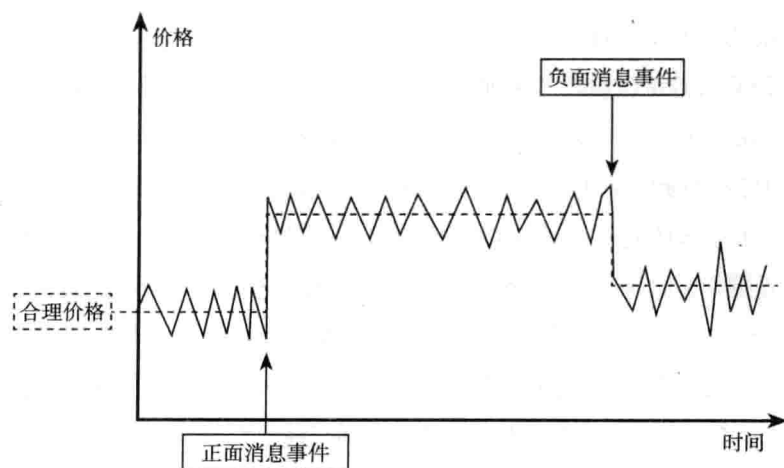


图 7-1 有效市场对于正面和负面消息的反映

市场有效性结果：好与坏

市场的有效性兼具好与坏两个方面。良好的市场有效性对于整体经济来说是向好的，但是糟糕的市场有效性对于技术分析而言，就是非常糟糕的。之所以说良好的市场有效性对整体经济是向好的，是因为理性价格为资产价值发出了重要信号；反过来，它会促进诸如资本与劳动力等稀缺资源的有效配置。这些因素会促进经济的发展。¹⁰

然而，有效市场中的价格把所有无效技术分析法则的形态以不可预知的形式呈现出来。保罗·萨缪尔森（Paul Samuelson）认为，金融市场价格之所以会遵循随机游走，是因为很多聪明的/理性的投资者行动的缘故。为了追求财富的最大化，他们买入被低估的资产并将其价格推高，然后再将这些被高估的资产卖出并把其价格打压得更低。¹¹ 如果我们把这种逻辑极端化，也就是说让一只被蒙住眼睛的猴子通过投掷飞镖的方法在报纸的金融版面上选择股票的效果，与专家们经过认真挑选的股票的的效果是一样的。¹²

有效市场假说在可得信息具有超过均衡预期利润或收益率的预期利润和收益率的基础之上，排除了交易系统的可能性。¹³ 用通俗易懂的话来说，普通的投资者——不论是个人的，还是养老基金，或者是共同基金，都没有指望他们能够持续跑赢大市，而这些投资者对大量的证券资源进行的分析、挑选以及交易都是徒劳的。¹⁴ 根据有效市场假说，即使存在能够按照投资策略获得正收益率的可能，但是当这些收益率经过风险的调整后，他们的收益率将不会比买入并持有市场指数投资组合的策略高多少。

有效市场假说也声称，当市场中没有新的信息出现时，价格都会处于随机的振动状态，并且会以无偏差的方式围绕着理性价格水平的上下运行（见图 7-2）。由于这条水平线本身就具有不确定性，因此当价格在其上下运行时，也就没有可靠的技术性或者基本的指标以供参考。这也就是说，在有效市场中，账面价值低的股票和大家所熟知的基本指标，是不可能涨过账面价值高的股票的。这对于任何寻找边界幅度的人来说无疑是冷酷的。

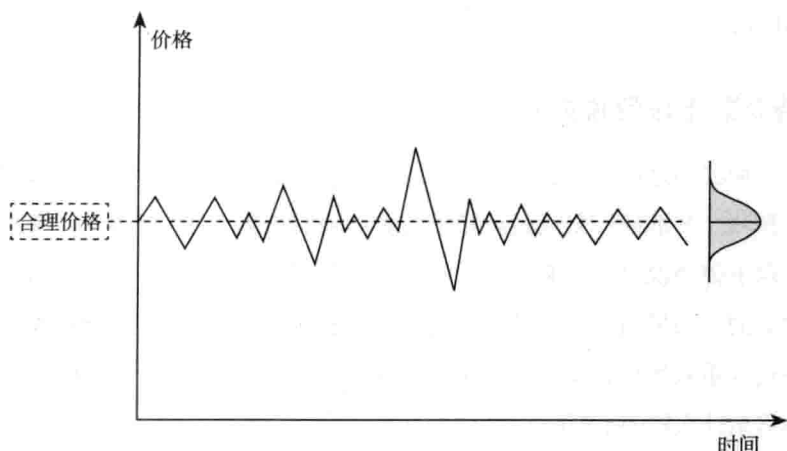


图 7-2 有效市场假说 (随机的和无偏差价格误差)

市场有效性的错误概念

对于市场有效性来说,始终存在着被普遍持有,但又错误的概念。

第一个错误概念是有效性要求市场价格必须随时和理性价格相等。实际上,市场有效性就是要求价格与理性价格通过无偏差的方式相背离。因此,与理性价值出现的正背离应该与其负背离相等,并且在任何时点上都应该具有资产价格在理性价格之上或者之下的同等概率。另外,这些背离都是随机的。因此,有效市场中价格偏差的时间序列就是一个随机变量,其概率分布的中心位于理性价格的位置,并且是对称的。

第二个错误概念是市场有效性暗示了没有人能够跑赢大盘。实际上,在我们已知的时间周期中,有 50% 的投资者可以跑赢基准指数,而剩下的 50% 的投资者的表现则弱于大盘。所以,没有单独的投资者或者基金经理可以长期跑赢大盘的概念也是错误的。对于数量众多的参与者来说,即使他们的策略完全没有优势,也会有一小部分的投资者可以获得持续跑赢大盘的收益率。纳西姆·塔勒布 (Nassim Taleb)¹⁵ 指出,即使我们假设基金经理和投掷飞镖¹⁶ (跑赢市场的概率等于 0.5) 的猴子的表现是一样的,并且经理人这个世界的数量又足够庞大,比如 10 000 个,那么在 5 年以后,仍然有 312 家公司的经理可以做到连续 5 年跑赢大盘。这些幸运的人就是那些实现成功登门拜访的推销员,而剩下的 9 682 家则还没有出现

在你的门前。

支持有效市场假说的实证

正如我们之前讨论过的，科学的假设起到了双重的作用：解释概念和进行预测。其准确性是通过将预测与新的观察结果相比较得出的。我们在第3章中曾经说过，只有在提供具体的歪曲实证的情况下，这种预测才是有意义的。如果新的观察结果与假设的预测相矛盾，那么这个假说要么重新表述，重新测试，要么被否决。然而，当观察结果与预测相符时，我们又会得出什么样的推断呢？

这件事现在正处于争辩之中。大卫·休谟阵营中的哲学家主张从来没有足够的确定性实证来证明一个理论，但是在实践科学的世界中，我们要充分证明一个假说。如果无数次的测试之后产生的观察结果与最初的预测一致，那么这对于科学家来说确实是意味着什么的。

所以有效市场假说就提议建立一个具有支持性实证的名单。由于有效市场假说具有两个可以检验的形态：半强式有效市场和弱式有效市场，¹⁷ 所以这些实证也具有两种不同的特点。半强式有效市场认为投资者是不可能利用公共信息来跑赢市场的。因为市场中包含了所有的基础数据和技术数据，所以半强式有效市场预测不论通过基本面分析还是技术分析，都是不可能打败市场的。这种非正式的预测并不足以证明它的可检验性，但是有一点是投资者们已经达成一致的，即投资者是不可能在波动的市场中持续获利的，我们为此进行了多次的尝试。

半强式有效市场也能够做出一些可以进行检验的具体预测。其中之一是证券的价格将会对影响价格的任何新信息做出快速且准确的反应。有效市场对信息做出适当反应的情况如图 7-3 所示。

预测显示了价格既不会对信息做出过度反应，也不会反应不足。这一点可以通过下面的方法进行检验。假设有一段时间价格的运行并不是很正常，并且系统性地对新信息做出了过度或者不足的反应。如果真是这样的话，就说明价格的运动是可以预测的。为了理解其中的缘由，我们来分析一下如果股票对信息产生过度反应（对好消息的反应是价格的上涨，对坏

消息的反应是价格的下跌)时的表现。由于有效市场假说主张价格必须紧密地围绕在理性的价格周围,那么在超过理性价格时,必然会导致价格向正确的理性价格回归。这种修正并不是随机的漫步,而是一种有目的的运动,就好像价格本身具有魔力一样画出一条理性的水平线。从某种程度上来说,这种运动的形式是具有预测力的。因为有效市场假说排斥可预测性,所以它也要排斥对消息产生过度反应的可能性。对牛市新闻假定的过度反应,以及随后系统性地向理性水平回归的运动如图 7-4 所示。

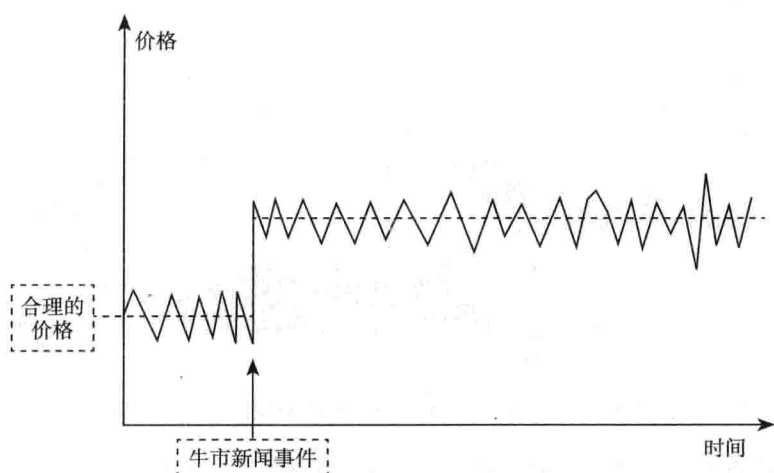


图 7-3 有效市场假说对牛市新闻的反应

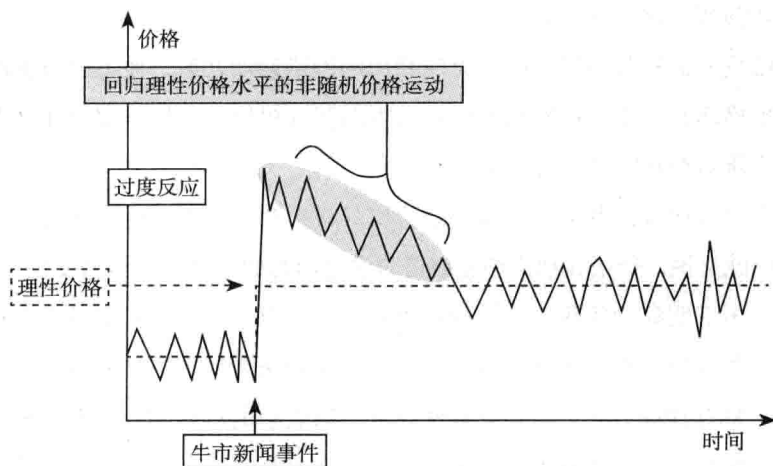


图 7-4 对牛市新闻(暗示了非随机价格运动)的过度反应

根据同样的逻辑，当价格的运动对经过证明的消息没有做出应有的反应时，有效市场假说也拒绝承认对消息做出的不足反应的可能性。在这种情况下，当价格继续向着新的理性水平延伸（波动）时，系统的价格运动也会随之变化。对牛市和熊市新闻事件的反应不足，以及由此而出现的向着理性价值进行的非随机运动（见图 7-5）。

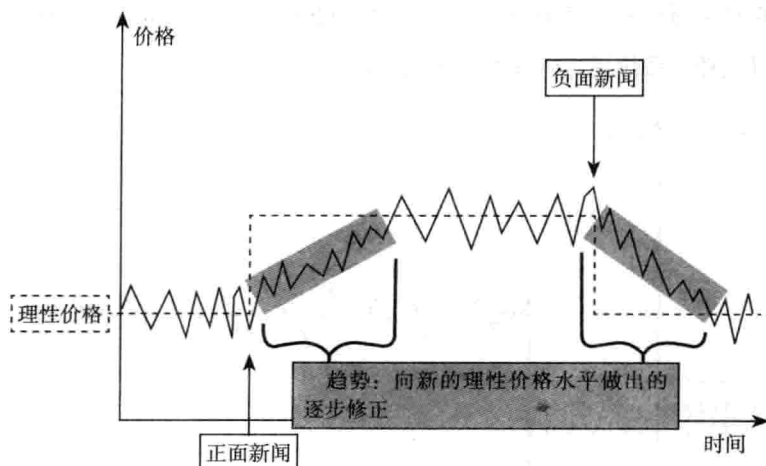


图 7-5 对可以解释非随机趋势的新闻做出的不足反应

受到实证支持的有效市场假说的预测认为，对消息的过度反应与反应不足在尤金·法玛的事件研究中是不会发生的。¹⁸他分析了所有类型公司的最新新闻事件，如收益和红利公告、吞并、并购等，发现在经历了由新闻所触发的最初几乎是瞬间且无法利用的价格运动之后，便不会再出现额外的价格运动。这些观察结果与有效市场假说提出的市场能够快速并准确地反应新闻事件的说法是一致的。¹⁹

有效市场假说的第二个预测是价格只有在消息出现的时候才会发生变化。也就是说，价格在缺少消息或者对没有内容的信息事件做出反应时，²⁰是不会出现明显的变化的。这两种预测提出了对于技术分析来说尤为重要的第三种预测。有效市场假说指出，在大众投资者当中出现的过时信息应该是不具有预测力的，而且其在使投资获取收益的过程中是没有任何价值的。很明显，所有被技术分析证明过的信息都是过时的，因此，有效市场假说预测所有的技术分析方法都是无效的。如果这一点能够得到证实的话，

那么技术分析的末日就要到来了。

所有用过时信息制定的策略应该是不具有预测力的这一假设在没有受到检验之前看上去没有任何意义。回想一下，有效市场假说认为基于过时信息制定的投资策略，在进行风险调整之后，是不可能产生超额利润的。以过时信息制定的策略获得收益的这一事实还不足以反驳有效市场假说。与有效市场假说相矛盾的是，它必须证明获得的超额收益是经过策略风险的调整之后得到的。

对经风险调整后的收益率为基础的投资策略进行判断是非常合理的。例如，如果某策略的年收益率是 20%，基准是 10%，但是该策略却表明投资者要面对 3 倍的风险，那么进行了风险调整以后，该策略并没有跑赢指数。问题出现了：如何定义并量化风险呢？

量化风险需要风险模型。最广为人知的模型就是资本资产定价模型 (CAPM)，²¹ 它解释了根据单一风险因子，证券收益率的系统性差异以及证券的相对波动率。也就是说，股票的风险是通过它相对于整个市场的波动率得到量化的。因此，如果一只股票的波动率是市场指数的两倍，它的收益率就应当是市场指数收益率的两倍才对。当这只股票获得了高于市场两倍收益率的时候，我们就把它称之为超额收益率。这种方法只有在衡量单一股票，或者持有多头股票头寸的投资组合的风险时，才能有效地运行。然而，CAPM 却不能把更多复杂策略的风险敞口进行量化。为了应对这种局限性，资本市场理论又提出了其他的风险模型。其中之一就是著名的套利定价理论 (arbitrage pricing theory, APT)，该理论把风险归结为许多独立的风险因子。²²

然而，有效市场假说的拥护者却给予了自己在任何时刻都可以虚构新的风险因子的权利，这使得他几乎无法对由过时信息制定的策略是不可能跑赢大盘的这一主张进行反驳。但是他们却能够以提出新模型的方式为一种策略获得的超额收益率是对风险的补偿的说法提供辩解。正如我们在第 3 章里解释的那样，如果风险因子在用过时信息制定的策略获得超额收益率被发现之前就确定了风险因子，那么这就不是科学的问题了。然而，在跑赢大市的策略被发现之后才提出新的风险因子（多个风险因子）无非就

是一个临时设定的假说。这是为了让一种理论达到免疫的目的这一事实之后做出的虚假解释，这对于有效市场假说来讲就是弄虚作假的表现。这种做法在以寻求充分的理由为己任的科学领域是不受欢迎的。而且，我们在第3章中曾经讨论过，科学理论的信息等同于它为证伪提供了机会。对任何一个或者所有自相矛盾的实证进行不受约束的辩解，实际上都会消耗该理论的信息内容。因此，试图随意编造新的风险因子来解释每一个用过时信息制定的跑赢大市的策略，实际上是将有效市场假说变成了无效理论。

能够证明弱式有效市场的实证表明，从过时的价格和指标中获得的价格运动等信息都是没有价值的，这一理论是以对自相关进行的研究为基础的。它衡量的是在各种滞后区间²³中，价格的变动与之前价格产生的变动之间的线性关系程度。这些研究证明了价格的变动（如收益率）实际上是线性无关的。这就意味着当前和过去收益率的线性函数是不能够被用于对未来的收益率进行预测的。²⁴由于这些发现，使有效市场假说的支持者得出了结论，即证券的收益率是不可预测的随机游走。

然而，从某种程度上来说，自相关研究是相对较弱的测试，因为它们只能检验出线性相关。如果时间序列只表现为非随机的方式，那么寻找线性结构就是确定线性相关的唯一方法了。埃德加·彼得斯（Edgar Peters）²⁵、罗闻全和 A. 克雷格·麦金勒（A. Craig MacKinlay）²⁶指出，可以替代的数据分析方法显示出金融市场的时间序列并不是呈随机游走形态的。例如，罗闻全和 A. 克雷格·麦金勒使用被称为“方差率”的检验统计量，它可以发现更加复杂（非线性）的非随机行为，而这种行为将会用于隐形的自相关研究。

挑战有效市场假说

优秀的理论在两个方面是一致的，即具有能够自圆其说的逻辑以及与观察的实证相一致。因此，一种理论既可以接受证明其自身矛盾的挑战，又可以接受实证的反驳。

如果某一理论证明或者预测到某些与该理论相矛盾的事情，那么该理

论就是自相矛盾的。我们在本章前面的内容对此列举了实例。我认为它是流行的技术分析的基本前提，即技术分析认为价格是对所有信息的完全反映，并且否认了技术分析的另外一个核心假设——价格包含了具有预测性的信息。

好的悖论与差的悖论

市场有效性的含义之一是，常识渊博的投资者获得的收益率应当比那些知识水平低下的投资者所获得的收益率低。这种说法是遵循着市场有效性意味着无论是已知的还是可知的信息都已经在证券的价格中得到了反应这一假设的。在有效市场中，是不存在对精明的投资者来说是竞争优势，或者说是表现差劲的投资者来说是竞争劣势的情况的。

然而，这种含义与有效市场假说的另外一种假设是相矛盾的，即理性投资者的套利行为是可以促使价格向理性水平运动的（参见本章后面的题为“有效市场假说的假设”一文）。只有在套利者拥有比不理性的、差劲的投资者更多的资金，并且具有更强的价格推动力时，套利者才能够作为理性价格的强制执行者来做。握有更多的交易资金意味着精明的套利者在过去必须获得比消息不灵通（差劲）的投资者更高的收益率。要么市场的价格由精明的、富有的参与者制定，要么不是。有效市场假说则包含了这两种含义，这难道不是矛盾的吗？

信息矛盾的代价

另外一个与有效市场假说逻辑不一致的就是信息成本。我们认为，对为获得消息而付出的时间、金钱、智慧等成本，以及把这些消息转化为有效的投资策略的行为进行假设是合情合理的。同时，有效市场假说声称这种消息对付出以上各种成本的投资者来说，是不可能赚取增值型收益的。回想一下，有效市场假说主张信息会马上通过价格反映出来。这一含义表明，无论做多么广泛或有深度的研究，其结果都将是毫无价值的。

然而，如果有效市场假说对信息的收集和处理没有得出任何结果的话，那么投资者就没有动力为这些消息而付出任何成本了。在没有人发现消息

并利用消息的情况下，这些消息就不可能被价格所体现。换句话说，如果为了获得有效的市场信息而付出高昂的代价，那么只有在用经过调整的超额收益率作为补偿的情况下，投资者才会有动力去做这件事。因此，矛盾就产生了，即有效市场假说要求信息的寻找者必须为他们的努力得到补偿，并且同时对他们将要做的事情予以排斥。

这种矛盾在格罗斯曼和斯蒂格利茨发表的《论信息有效性市场的不可能性》这篇论文中得到了令人信服的论述。²⁷ 他们都认为低效率对于推动理性投资者参与信息收集和处理方面是不可或缺的。理性投资者从市场的定价误差中赚取的收益是对他们所有努力的补偿，而投资者获得的收益又起到了让价格向着理性水平方向运动的效果。理性投资者的利润是通过噪声交易者和流动性交易者的损失获得的。噪声交易者是指根据他们认为包含丰富信息（实际上根本就没有）的信号进行买卖来增加其现金储备的投资者。理性投资者同样会从通过交易获得现金储备量增加（流动性交易者）的投资者出现的损失当中赚取利润。如果被发现的时间信息与反映在价格之中的时间之间的关系出现滞后时，理性投资者便会从中赚取利润，但是，有效市场假说却排斥价格可以逐渐调整的可能性。这就是与有效市场假说逻辑另一个不一致的地方。

类似的逻辑也适用于对信息的处理（如佣金、下跌、买卖差价）。除非能够得到对其交易成本的补偿，否则投资者就没有进行交易的必要。根据有效市场假说的要求，投资者进行交易活动的目的是把价格推向理性价格水平。从本质上来说，任何限制交易的因素——交易成本、生产信息的成本，削弱卖空的法则等因素都会限制市场实现有效性的能力。我们所持有的对从事的交易是没有补偿的这种逻辑是不可取的。

有效市场假说的假设

为了理解有效市场假说的批评者提出的论证，我们有必要理解 3 种逐步走弱的假设：①投资者是理性的；②投资者对证券价格的定价偏差是随机的；③市场中通常会有理性的套利者捕捉到定价偏差。我们将对这 3 种假设逐一进行详解。

有效市场假说争论的第1个观点就是,投资者总体来说都是理性的。回想一下,理性投资者总体上都是希望正确定义证券的价格的。这意味着价格将会尽可能准确地反映某一证券未来现金流与风险特性的当前贴现值。理性投资者对新获得的信息反应迅速,当获得利好的消息时,他们会迅速推高证券价格;当得到利空消息时,就会打压证券的价格。结果,证券的价格会因为新的消息而被经常调整。²⁸

即使第1种假设是不正确的,有效市场假说仍然坚持认为价格会向着理性水平进行调整,这是源于该假说的第2种假设:单个投资者的估价错误是互不相关的。也就是说,如果一位投资者错误地把某一证券的价格估计的过高,那么这种错误就会被另一个错误地将证券价格估计过低的投资者所推翻。总的来说,这些估价错误会自我抵消,且平均价值为0。也就是说,定价错误是无偏差的。这种说法被认为是正确的,是因为容易出现差错的投资者有可能运用不同的,甚至是无效的投资策略。因此,他们的行为应该是不相关的。坦率地说,就是愚笨的投资者根据一个毫无信息内涵的信号买入证券,而另外一个愚笨的投资者同样根据这个毫无信息内涵的信号卖出证券。实际上,他们相互交易的结果使得证券的价格非常接近理性水平。

此外,就算第2种假设没有使价格保持在理性水平上,有效市场假说还可以运用第3种假设,即如果理性投资者做出的错误定价存在偏差(平均价值大于或者小于零),套利者便会在此时出手,他们会注意到价格已经与理性价格出现了背离。套利者会在价格过低的时候买入该证券,而在价格过高的时候将其卖出,通过此举强行让价格向理性价格水平回归。

在有效市场假说领域中,套利几乎等于免费的午餐。套利交易不会承担风险、没有资本金的要求,可以获得稳定的收益。正如有效市场假说的支持者预想的那样,套利交易可以被表述为:假设现在有两种金融资产——股票X和股票Y,我们把它们在同样的价格出售,且两只股票承担的风险也相同,然而,它们却具有截然不同的未来预期收益率。显然,其中的一只股票出现了定价错误。如果资产X具有更高的未来收益率,那么借助定价错误的优势,套利者可以在买入资产X的同时卖空资产Y。随着

具有相同目的套利者的活动的进行，每一只股票的价格将会向其正确的基本价值²⁹靠拢。通过这种方式获得的市场有效性说明，总有一些套利者已经准备好，愿意并且能够抓住这些机会。越富有的人越有把价格推动到正常水平的能力，而随着非理性投资者的陆续破产，他们也会失去将价格推动到远离均衡水平的能力。正如弱小的物种会被生态系统所淘汰一样，那些无知的、非理性的投资者也将会被排除在市场之外，而理性的、学识渊博的投资者才能留下来。

有效市场假说所做假设的缺陷

这部分分析的是有效市场假说所做的每一个假设，以及它们是如何犯错误的。

1. 投资者是理性的

投资者并没有表现出像有效市场假说预测的那样理性，很多投资者都会受到毫不相关的消息的影响，经济学家费希尔·布莱克（Fischer Black）把这种现象称为“噪声信号”。³⁰ 尽管这些投资者认为他们的做法是非常明智的，但是他们所获得的收益与那些根据掷硬币决定买卖的投资者获得的收益是相同的。追随噪声信号的投资者就是噪声交易者的代表。

事实上，投资者对于自己与理性背离感到内疚。例如，他们没有把自己的投资多样化，并且积极地进行交易，他们通过卖出上涨的股票并牢牢持有亏损的头寸反而增加了自己的税负压力，他们需要为交易共同基金支付高昂的费用。另外，这样做的结果就是亏损的恶性循环的开始。投资者与经济理性最大化的背离证明是普遍存在的和系统性的。³¹ 我们把投资者与理性背离的情况分为3种：不精确的风险评估、差劲的概率判断以及非理性的决策框架。我将依次对这3种情况进行讨论。

纽曼（Neumann）和摩根斯特恩（Morgenstern）提出，投资者不会按照规范的观点对风险进行评估的，我们称其为“预期效益理论”。在完全理性的市场中，个人在面对众多的选项时，往往会选择那个具有最高预期值的选项。特殊选项的预期值等于一组乘积的总和，即每个乘积等于每一个可能结果的概率乘以这个结果对于决策制定者的价值或者效用。但问题是

人们通常并没有这样做。研究显示,当投资者进行风险评估和做出决策时,会产生大量的系统性错误。

这些错误属于一种被称为“预期理论”的框架,它是由卡尼曼和特沃斯基³²于1979年提出的。该理论解释了人们是如何在不确定的条件下进行决策的。例如,投资者通过持有亏损的股票而卖出可以盈利的头寸来避免实际亏损带来的精神痛苦。它同时又解释了投资者为什么高估了高风险投机活动出现的概率会很低这一现象。

投资者与有效市场假说的理性观点相背离的第二种情况就是他们对概率的判断。有效市场假说假定,当投资者获得新的消息时,他们会依照贝叶斯定理对概率的评估进行更新,即通过理论上正确的方式把很多的概率合并成为一个公式。然而,遗憾的是,投资者并没有按照这种方法去做。我们在第2章中讨论过的一个普遍性的错误,即从小规模的数据样本中推导出宏大结论的小数错误。它解释了投资者会轻率地做出只根据单个季度盈利的短期结果来判断股票后市会快速上涨的原因。投资者对这种正收益率变化的小规模样本做出的过度反应会导致股票的价格被高估。

最后,投资者的决策还会受到如何描述(构架)选择的强烈影响。由于对决策进行的不准确构架,投资者会对其选择的真实预期价值产生误解。例如,当以潜在获利为基础对选择进行构架时,投资者会选择那个最有可能获利的选项,而至于获利的规模大小则无关紧要。然而,当以对潜在的亏损为基础对同样的选择进行构架时,投资者会假设以损失非常严重的风险来避免小规模亏损的确定性。这种错误被称为“处置效应”,它可以用来解释投资者会在即使获得的利润只有那么一点点的情况下,把盈利的股票快速地卖出,反而会在手中的亏损额出现大幅增加的情况下,却坚持持有亏损股票的原因了。

2. 投资者的错误是不相关的

通过心理学研究表明的投资者的错误是不相关的这一假设是矛盾的,它显示了在以不确定性为特点的情况下,投资者是不会随机偏离理性的。实际上,大多数人在这种情况下都会犯类似的错误,因此,他们违背理性的原则是存在相互关系的。很多投资者都倾向于买入同样的股票,因为对

股票进行直观推断法则的判断的需求使得该股票看上去有在未来升值的可能。例如，当很多的投资者发现某家公司具有三个季度的正收益时，他们就会认为高市盈率（P/E）的公司是成长型股票。在社会互助、羊群效应以及投资者之间，这个问题就显得非常严重。他们听信谣言，往往跟随着相互之间的错误行事，并且模仿交易同行的行为。³³

资金经理应该比投资者做得更好，但是他们也在犯同样的错误。他们为创建的投资组合会如此靠近基准而感到羞愧，而这个基准是用来衡量他们尽一切努力试图减少自己表现不佳的基准。专业的资金经理也会模仿这种行为，为了避免落在别人的后面，他们也会把这种从众行为带入到同样的股票当中。资金经理通过把表现上佳的股票添加到投资组合中来，并且把业绩糟糕的股票剔除出组合的方式来粉饰其投资组合，因此，在年底的投资组合报告中就会显示出资金全部投到了那些表现最佳的股票上。

3. 套利迫使价格向理性水平回归

这个假设之所以是错误的，是因为套利行为对将价格推向理性水平的能力是有限的。

首先，当证券出现价格错位时，是没有人会告诉你的。理性价格是未来现金流的贴现值，当然，这个值是不确定的。投资者试图估算未来的收益，但是他们的预测往往带有很大的误差。

其次，套利者对反转的价格走势没有足够的忍耐力，例如，当在偏低的价位买入的股票继续走低，或者已经抛售的股票仍然上涨的时候。在价格重回理性价格的水平之前，噪声交易者的行动可以轻易地把价格推到离理性价格很远的水平上。因此，即使套利者发现了一只真正价格错位的证券，那么在错误的价格被消除之前，价格错误也许会越拉越大。可能当反转的价格运动离理性价格越来越远时，套利者就不得不以亏损来平仓了。如果这种情况经常出现的话，那么噪声交易者就可以把理性投资者逐出市场。这就是我们所指的噪声交易者风险，³⁴同时它也限制了套利交易者的积极性和其所承担的义务。行为金融学家安德烈·施莱弗（Andre Shleifer）声称，虽然从表面上看套利交易是近乎完美的，但实际还是有很大的风险的，所以依然会对那些追逐套利交易的投资者形成限制。从本质

上来讲, 套利行为不可能总能实现合理定价。

噪声交易者风险是美国长期资本管理公司 (LTCM, 对冲基金公司) 在 1998 年秋天破产的主要原因。该基金的破产几乎波及了整个金融市场。最终, 由 LTCM 确定的定价错误得到了修正, 但是由于他们对头寸的过度杠杆化, 使得该基金失去了在短时间内持有更大定价错误头寸的持久力。

对杠杆的非正确使用是另外一个削弱套利在推动市场有效性中发挥作用的因素。因此, 即使套利者能够准确地确定被高 / 低估的证券, 假如他们使用太多的杠杆, 他们也会被市场所淘汰。对于有利的选择来说是没有合适的杠杆可以用的 (如对积极的估计进行投机活动)。³⁵ 如果过度地使用杠杆, 那么即使是积极的预期, 其破产的概率也会增加。

这是在一堂被我称为“赌场之夜”的课上我对我的学生们说的。在课上, 我们做了一个游戏, 假设我们每个人在开始的时候都拥有 100 美元, 我的一个朋友以投掷硬币者的身份加入进来, 而学生们则要决定为下一次掷硬币所下的赌注。如果出现的是头像那面, 那么他们所得到的报酬就是赌注总额的 2 倍; 如果出现了背面, 仅仅会损失这次所下的赌注。我们把这个游戏玩了 75 次。这个游戏取得了非常好的预期效果。³⁶ 然而, 即使是在这个充满良好预期的游戏中, 很多学生却因为自己太过激进而以输光告终 (对每次下注投入过多的资金)。20 世纪 50 年代, 由贝尔实验室的电器工程师凯利 (Kelly) 发明的凯利公式 (Kelly Criterion) 说明了在每次掷硬币时应该下注多少“最佳”的金额才能够让下注者手中的资金增长率达到最大化。这个数字取决于赢的概率和平均获胜占平均失败的比率。在这个特殊的游戏当中, 在每次掷硬币时下注的理想值是 0.25。如果超出了这个水平, 下注者将会面临更大的风险, 而不是快速增长利润。如果有人把下注的数值定为 0.58, 即使是面对有利的预期, 那么他也很有可能输光所有的家当。以上就是套利者运用太大的杠杆处理利好消息时将会发生的情况。

另外一个制约套利者促进有效价格能力的因素是缺少完美的替代证券, 理想的 (无风险的) 套利交易是指对两种具有完全相同的未来现金流和风险特征的证券做到同时买进和卖出。任何在不包含这两个特征的基础上进

行的套利交易都是有风险的。这种风险限制了套利活动能够推动价格走向有效水平的程度。在广泛的资产类别中，比如说在 2000 年春天，所有的股票都出现高估的情况下，就没有一种可以替代的证券扮演多头的角色，来对针对整个股票市场的做空交易进行对冲。即使是在套利交易中的证券也会有类似的替代品这个案例中，也存在着一定的风险。每只股票都有自己独特的或者是非系统性风险，这就使得套利者不得不加倍小心。例如，如果通用汽车公司的价格被低估，那么就通过大量买入福特公司的股票进行抵消，但是这里所面临的风险就是针对福特公司的牛市事件或者针对通用汽车公司的熊市事件有可能对套利者造成亏损。

目前还存在着另外一种由有效市场假说假设的限制套利发挥价格政策作用的因素。套利者没有永远也用不完的资金来对定价错误进行修正，也没有绝对的自由去捕捉任何一个或者是所有的机会。大多数进行套利的投资者利用对冲基金的形式为其他投资者管理资金。在管理合约中对基金经理操作行为的程度予以限制。如果没有完全自由的手和不受限制的资金，有些定价错误仍然是无法套利的。

由于这些限制的结果，有效市场假说关于所有套利机会将会被消除的假设就显得过于简单化了。

挑战有效市场假说的实证

有效市场假说不仅在逻辑的不一致方面漏洞百出，而且其预测也与大量的实证相矛盾。下面这部分内容将对这些发现进行总结。

1. 价格的极度波动

如果证券的价格像有效市场假说声称的那样是与基本价格紧密相连的，那么价格波动的幅度就应当与潜在的基本价格的变动幅度相似，然而，研究显示价格要比基本价格的波动大得多。例如，当我们把基本价格定义为未来红利³⁷的净现值的时候，基本价格的变化是无法解释价格的大幅波动的。2000 年春天科技股泡沫的破裂和 20 世纪 80 年代末日本股票价格出现的泡沫就是价格运动无法通过基本价格的变动进行解释的佐证。

有效市场假说还预测价格的大幅变动只有在重要新消息进入市场的时候才会出现，而实证却并不同意这个观点。例如，1987年股票市场的崩盘就没有和证明价格会下跌超过20%的消息同步。1991年的一项由卡特勒（Culter）研究，施莱弗³⁸所援引的研究对自第二次世界大战结束以来50个最大的单日价格运动进行了分析。其中大多数的价格运动与当日公布的重要消息无关。在另一个类似的研究中，罗尔（Roll）³⁹表示，冻橘子汁的未来价格波动通常与影响市场的天气因素无关。罗尔还指出，单个股票的运动通常与公司的消息无关。

2. 关于使用过时信息进行价格预测的实证

对于有效市场假说来说最确凿的证据就是某项研究显示，根据公众已经知道的（过时的）信息可以预测到有意义的价格运动。换句话说，以过时信息为基础制定的策略可以产生经过风险调整且跑赢大市的收益率。如果有效市场假说主张的价格可以迅速地包含所有的已知信息的说法是真实的，那么这就是根本不可能出现的。

3. 横截面的可预测性研究是如何工作的

2001年3月发行的《金融杂志》指出，许多预测性的研究表明，那些能够让公众获得的信息并不能完全通过股票价格反映出来，而且很多基于这些信息的策略都是可以盈利的。⁴⁰ 这些研究衡量的是以公共信息为基础的指标能够预测股票相对绩效的程度。⁴¹ 指标测试的内容包括市盈率、市净率、股票近期的相对价格绩效等。

可盈利性研究使用的是横截面数据。也就是说，这项研究对在某一给定时间点的大量股票的横截面进行分析，例如，截止到1999年12月31日，在标准普尔500指数中的所有股票。在这一时点上，所有的股票根据某一种正在对其预测力进行分析的指标进行排序，⁴² 比如说股票在过去6个月（价格势头）的收益率。在此之后，这些股票在排序的基础上被分为若干个投资组合，且每一个投资组合中包含数量相等的股票。10个投资组合为一个常数。因此投资组合1包括所有在过去6个月中收益率排名前10位的股票。也就是说，这个投资组合1当中包括的是在过去6个月中收益

率高于其余 90% 股票的那 10 只股票。第 2 个投资组合是由排名第 80 到第 89 位的股票组成, 依此类推, 直到最后一个由过去 6 个月中绩效最差的 10 只股票组成为止。

为了确定这 6 个月的趋势是否具有预测力, 将排名最靠前的投资组合的未来绩效 (1) 与排名最后的投资组合的未来绩效 (10) 进行比较。向前的预测水平线是一个时期, 因此, 如果利用月度数据, 就是指一个月。⁴³ 这些研究把对指标的预测力进行的量化作为多头投资组合与空头投资组合的比较。换句话说, 假设投资组合 1 中包含的股票都是多头, 而投资组合 10 当中包含的股票全都是空头。例如, 如果在一段给定的时间内, 多头投资组合的收益率是 7%, 而空头投资组合则亏损 4 个百分点 (如卖空股票上升了 4%), 那么多头 VS. 空头的策略就会获得 3% 的利润。尽管在这个案例中用于排列股票所使用的指标是 6 个月的价格趋势, 但是在投资组合被建立之时已知的任何信息都是可以利用的。这一概念如图 7-6 所示, 其中所使用的指标是每只股票的市盈率。

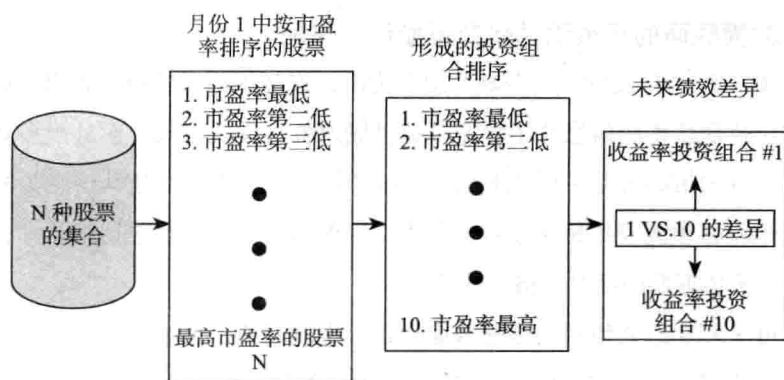


图 7-6 确定市盈率的预测力的横截面研究

我刚才描述的是如何对单一月份 (1999 年 12 月) 进行横截面研究的。然而, 横截面研究也可以对多个月份 (横截面的时间序列) 进行研究。因此, 衡量某一指标的预测力等于在一段经过扩展的时间段中, 用投资组合 1 的平均收益率减去投资组合 10 的平均收益率。这种研究通常分析的是若干候选的预测变量 (见图 7-7)。

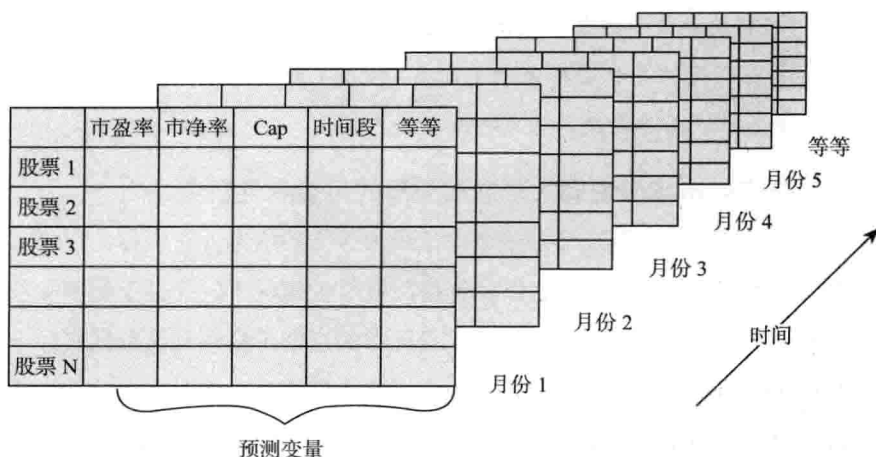


图 7-7 横截面时间序列研究

4. 与半强式有效市场假说相矛盾的预测性研究

有效市场假说的半强式形式是有效市场假说最为大胆的可测试版本。⁴⁴ 该形式主张, 存在于公共领域内的消息、基本面消息或技术性消息都不可能产生超过经风险调整的市场指数的收益率。大量表现良好的横截面时间序列研究的本质是: 利用已经过时的公共信息, 价格趋势在某种程度上是可以预测的, 并且可能获得高于经过风险调整的收益率。在此, 我把一些重要的研究成果加以总结:

- **小盘效应:** 一只股票的全部市值是通过已经公开发行的股份乘以其每股的价格计算得出的, 它是对未来收益率的预测。⁴⁵ 在投资组合里的股票中, 最低市值的投资组合中的股票可以获得每年约 9% 的收益率, 这比排名最高的投资组合中的股票获得的收益率要多。⁴⁶ 这种效应在 1 月尤为显著, 而最近的研究表明, 这种指标的预测力自 20 世纪 80 年代中后期开始便消失了。
- **市盈率效应:** 低市盈率的股票要比高市盈率的股票表现突出。⁴⁷
- **市价与账面价比率效应:** 低市价与账面价比率的股票要比高市价与账面价比率的股票表现突出。⁴⁸ 根据账面价值与股价比率, 价格越便宜的股票, 其每年的表现几乎超过价格最贵的股票的 20%。
- **利用技术信息获得的异常收益:** 公布非预期收益或者给出业绩指引

(根据消息公布后第二天出现的强烈的量价配合所决定)的股票在下一个个月获得的年收益率差额(多头-空头)超过了30%。⁴⁹这一策略是基本面信息与技术信息相结合的产物。

5. 与有效市场假说弱式形态相矛盾的可盈利性研究

有效市场假说的弱式形态是该理论最为保守的版本。它声称,在只有公共信息的子集、过去的价格以及价格回报⁵⁰的情况下,无助于超额收益的获得。下面的研究显示了用于技术指标中的过期价格和其他数据实际上是很有帮助的。这对于技术分析来说无疑是个好消息,而对于有效市场假说而言则恰恰相反。这个最狭义、最缺乏自信以及最难造假的有效市场假说版本一直受到实证的反驳。

► **趋势的持续性:** 杰加迪西(Jegadeesh)和蒂特曼(Titman)⁵¹在1993年所做的测试表明,价格趋势在过去的6~12个月的时间里是持续的。换句话说,一只股票在6~12个月期间的强势表现具有在未来的6~12个月的时间里依然表现上佳的倾向。他们的研究对持有的多头策略——前6个月获得最高收益的股票(排名最高)和持有的空头策略——前6个月绩效最差的股票(排名最低)进行了模拟。随后,将多头和空头的投资组合持有6个月。这种策略获得的年收益率是10%。根据这组实证,即使是有效市场假说的领军人物之一的尤金·法玛,也不得不承认股票过去的收益率可以对未来的收益率进行预测。⁵²这对技术分析来说是一次巨大的胜利!

► **趋势的反转:** 尽管在过去6~12个月的时间里具有强烈趋势的股票仍然有可能在接下来的6~12个月里保持这种趋势,但是当对这种趋势进行更长时间跨度的测试时却会出现某些差异。对该强烈趋势今后3~5年的走势进行的测试表明其具有出现反转形态的倾向。德邦特(De Bondt)和塞勒(Thaler)对一种策略进行了测试,即买入过去5年走势最差(亏损)的股票,并卖出过去5年走势最好(盈利)的股票,在接下来的3年中,⁵³多头与空头投资组合相比得出的是平均8%的收益率。最重要的是,收益率的差异并不能

归因于风险。之前的失败者（未来的胜利者）比之前的胜利者（未来的失败者）承担的风险要低。⁵⁴ 这一点反驳了有效市场假说的高风险产生高收益这一核心观点。此外，它还支持最简单的技术分析形态就是最有效的这一观点。

► **非反转趋势：**当用接近其 52 周高点的方法，而不是用其之前的收益率对股票的趋势进行衡量时，利润会更多，并且不会出现趋势反转。⁵⁵ 该项研究的作者推测投资者对之前的 52 周高点已经形成了固定心理。众所周知，锚定可以防止人们对新信息进行适当的调整。这项研究的作者迈克·库珀推测，这会防止接近其 52 周高点的股票对新的 important 发展趋势做出它应有的反应。延迟的消息回复产生了纠正错误价格的系统性价格趋势（势头）。

► **通过交易量确定趋势：**来源于进一步支持技术分析有效性的研究表明，当交易量与价格趋势共同使用时，即使是更高的收益率也是可以获得的。也就是说，协同效应可以通过价格与交易量指标的结合来获得。这一结合体的收益率在 2% ~ 7%，这远远超过了仅仅⁵⁶ 使用价格趋势获得的收益率。交易量大且具有积极价格趋势的股票接下来的表现要强于只有积极价格趋势的股票；交易量大且具有消极价格趋势的股票接下来的表现要低于只有负价格趋势的股票。换成另一种说法就是，交易量大的股票具有更强的趋势识别力。

为有效市场假说寻求保护

当一个已经确立的理论与实证相矛盾的时候，它的支持者不会轻易把它推翻。对经过训练的科学家们来说，坚持信念的认知机制（见第 2 章）使得他们不会像普通人那样去做事情。挑剔的科学观察者们认为，那些追随歪曲理论的人，是永远也不会改变他们的思想的。他们在有生之年已经为此投入了太多精力，而随着时间的推移，他们就会逐渐消失。

40 年来，有效市场假说作为金融学的基础表现得越来越好，它的支持者还没有打算低头认错。有效市场假说的两位支持者——尤金·法玛和肯尼斯·弗兰奇曾经说过，通过诸如市账率和市值等过时信息制定的策略获

得的收益无非是对风险的合理补偿。回想一下，有效市场假说不会否认利用过时的公共信息可以获得利润的可能性，这仅仅是说当这些收益进行风险调整时，它们的表现还不如投资指数基金的表现。

只要“风险”一词还没有被定义，有效市场假说的拥护者就会在事实之后虚构出新的风险形态。正如艾略特波浪理论所证明的那样，事后（after-the-fact）允许对以前的观察结果进行解释或者辩解。而在我看来，这就是尤金·法玛和肯尼斯·弗兰奇所做的事情。⁵⁷他们发明了一种使用三种风险因子的特殊风险模型来取代以前的只使用一种风险因子（即股票相对于市场指数的波动性）的资本资产定价模型。⁵⁸更加方便的是，尤金·法玛和肯尼斯·弗兰奇决定加入的两个新风险因子是市净率和市值。通过引入这些作为风险的替代，法玛和弗兰奇恰当地为它们的预测力进行了辩解。根据新的风险模型，小型（低市值）企业或者是低市净率的企业基本上是风险更大的企业，因此，它们就必须拿出更高的平均收益率来补偿那些愿意投资于它们的投资者。相反，由于大市值企业的股票更安全，那么高市净率的股票实际上就是那些企业未来前景更加确定的公司的股票，它们的平均收益率之所以较低，是因为它们把投资者的投资风险降到了更低水平。⁵⁹

这只不过是為了挽救将要消亡的理论所做的一个特殊事实之后的解释罢了。如果在发现这些变量可能获得超额的收益率之前，法玛和弗兰奇已经预测到了市净率和市值是有效的风险因子的话，那么实际情况就会完全不同了。它可以通过连续的观察来确定有效市场假说的演绎结果，并表明有效市场假说是一种强有力的理论。然而，法玛和弗兰奇并没有这样做。他们在市净率和市值之后创造的新风险因子已经表明可以产生超额收益率了。

正如我们前面指出的那样，法玛和弗兰奇通过提出低市值和低市净率是企业具有高度破产风险的信号证明了他们的新风险模型。如果真是这样的话，就暗示（预测）了价值策略（买入低市净率的股票）和小市值策略在经济形势不好的时期里，会产生低于预期水平的收益率，而那些身处困境的企业是最容易遭受打击的。

实证与上面所做的预测是矛盾的。一项 1994 年的研究发现，没有证据可以证明价值策略在经济环境不好的时期表现得更差。⁶⁰而且，在过去的 15 年中，小市值公司的股票获得超额收益率现象的消失又提出了一个问题，如果市值规模确实是一个合理的风险因子，那么基于该策略的收益率就应当继续获得风险溢价。最终，不论是法玛－弗兰奇的风险模型，还是其他有效市场假说模型，都不能解释价格动力指标的预测力，或者是价格动力与交易量产生的共同影响。

行为金融学：随机价格波动理论

当现存的理论面对矛盾的实证时，对新理论的需求就会出现。有效市场假说现在连续提出与其预测相悖的实证。现在是新理论登场的时候了。

行为金融学中这一相对较新的领域已经提出了各种有关金融市场行为新兴理论的变体，这些理论对有效市场假说所不能解释的现象进行了解释。这些理论之所以具有科学意义，是因为它们并不是简单的解释，而是对经过连续观察研究所确定的检验做出了预测。

行为金融学包含认知心理学、经济学以及社会学等要素，并对投资者会与理性相背离以及由此导致市场脱离完全有效性的原因进行了解释。在考虑了情绪、认知错误、非理性推断以及群体行为所产生的影响之后，行为金融学对过度的市场波动，以及用过时信息制定的策略能够产生超额收益率的问题给出了简单的解释。

行为金融学并不假设所有的投资者都是非理性的，相反，行为金融学把市场视为处于不同理性程度的决策者的混合体。当非理性的投资者（噪声交易者）与理性的投资者（套利者）进行交易时，市场可能会背离其有效的价格。实际上，市场的有效性不是一种缺少可能性的特殊状态，而是一种具有更多合理性的市场状态，在这种状态下，价格很有可能偏离理性水平，然后朝着理性水平做系统的可以预测的运动。⁶¹这也是以过时信息为基础制定的系统策略能够盈利的原因。

这一点对技术分析来讲无疑具有讽刺意味。正如我们在第 2 章中讨论

的那样，认知性错误可以解释人们是如何在缺少有力的支持性实证或者在面对矛盾的证据时，对主观技术分析的有效性形成错误的观点的。然而，认知性错误同时还能够让客观的技术分析方法起作用，并且对造成系统价格运动的市场非有效性的存在进行解释。以上内容不仅说明了主观技术分析实践者的愚蠢，而且还说明了一些主观技术分析方法的合理性。

行为金融学的基础

行为金融学是以两个支柱为基础的：套利对修正价格误差能力的限制和人类理性的限制。⁶² 当这两个概念结合在一起时，行为金融学可以预测具体的能够产生系统价格走势的市场效率的偏差。例如，在特定的条件下，可以预测趋势是保持不变的，而在其他条件下，可以预测趋势反转。我们将分别对这两个支柱进行分析。

1. 对套利的限制

我们先简要回顾一下，套利并不是有效市场假说假定的有效价格的完美执行者。缺少完美的证券替代品把套利从无需承担风险、没有资本金要求的交易转变成成为需要资本金并会招致风险的交易。即使现在有非常合适的替代证券，但是在回归其理性价格之前，仍然会存在未来的价格偏离理性价格的风险。此外，为套利者提供交易资金的投资者既没有耐心，也没有无限的资本。结果是，这些投资者根据被认为可以利用的机会类型，而不给套利者留有任何回旋的余地。

这些限制解释了证券的价格通常不能恰当地对新信息做出反应。有时会出现对价格的反应不足，有时还会出现对价格的反应过度。我们再补充上噪声交易者的影响这一点，噪声交易者是根据不包含信息的信号进行操作的，⁶³ 由此我们可以很清楚地看到，在较长的一段时间里，证券的价格是如何系统地偏离理性水平的。

2. 人类理性的限制

对套利的约束能够预测无效性的出现，但是它们不会单独预测在什么情况下这种无效性会出现。例如，对套利的约束不能告诉我们在哪种情况

下市场对新信息做出的滞后反应要强于过度反应。这也就是行为金融学的第二个支柱所说的，即人类理性的限制。

我们在第2章中学到的认知心理学已经表明，在不确定的情况下，人类的判断倾向于犯预测性（系统）的错误。考虑到人类判断的系统性错误，行为金融学可以预测背离市场有效性的类型，这种情况很有可能发生在特定的环境之下。

在行为金融学的两大支柱当中，更加便于理解的就是对套利的限制，而非投资者的非理性。这是因为套利者被期望是理性的，而经济理论则紧紧抓住理性参与者的行为，而不是非理性参与者的行为。例如，目前还尚不清楚在金融领域中，哪种具体的认知性偏差与系统性判断错误是最重要的，但是一幅图画正在逐渐形成，这部分内容描述的是，行为金融学对非理性投资者行为的理解，以及由此产生的系统价格运动的当前状态。

与传统的金融学理论相反的是，行为金融学主张投资者会做出具有偏差的判断和选择。实际上，市场价格之所以能够系统性地背离理性价格，是因为投资者系统性地背离了完全理性。然而，这些错误并不是无知的表现，而是人类为了应对复杂性和不确定性而对人类智力进行开发所采用的有效方式的结果。

行为金融学 and 有效市场假说具有相似性，是因为二者都主张市场会越来越向好，也就是说，价格最终会向理性的股价靠拢。它们之间的不同点也仅仅在于背离和持续时间：有效市场假说认为价格背离理性水平是随机的和暂时的；行为金融学则认为某些背离是系统性的，而其持续较长的时间也被一些投资策略所利用。

心理因素

在第2章中，我们看到了认知性错误影响主观的技术分析师的方式，以及由此导致的错误观念。在行为金融学的条件下，我们研究了认知性错误是如何影响投资者的，以及因此而导致的系统价格运动。我们在此将要讨论的认知性错误既不完全不同，也不是独立运转的。然而，为了清晰起见，它们都是单独表现的。

1. 保守主义偏差、确认性偏差以及信仰惯性

保守主义偏差⁶⁴是指忽视新信息的一种倾向。因此，人们在新信息得到证明的情况下，不能对他们以前的观念进行修正，之前的观念则更加趋于保守。正是因为这一点，投资者具有对与证券价值有关的新信息反应不足这种倾向性，所以证券的价格无法充分地对信息做出反应。然而，随着时间的推移，价格会系统性地向其价值靠拢。这种逐步的、带有目的性的调整会以价格趋势的形式通过图表表现出来，请参见图 7-5。

保守主义偏差得到了确认性偏差的支持，而确认性偏差导致人们接受与以前观念相一致的实证，并且拒绝相信与以前观念相矛盾的实证。因此，当投资者对某只证券持有固定的观点，并且新信息又可以对这种观点予以证实的时候，投资者就会对这些信息更加重视，并且对此反应过度；相反，当新信息与之前的观念相矛盾的时候，投资者就会对这些信息持怀疑的态度，并且对此反应不足。

确认偏差也会影响投资者的观念，不管这些观念是什么，都会随着时间的推移让投资者对此观念表现出过激行为。造成这种情况的原因如下：随着时间的推移，新闻是以确认性信息和矛盾性信息的随机结合体的形式出现的。然而，确认性偏差则导致了对这些信息处理方式的不同。确认性信息是已知的证明，因此它是对以前观念的加强，而矛盾性信息则不会被人们所轻易相信，因此它们对投资者之前的观念所产生的影响是微不足道的。结果是，随着时间的推移，以前的观念有可能被加强。

然而，如果一连串的矛盾性信息同时出现的话，那么以前的观念有可能会遭受彻底和非理性的弱化。即使这些普通的矛盾性信息是以随机序列出现的，投资者仍然会犯小数法则的错误，而且当众多与以前观念相矛盾的信息碎片连续出现时，投资者会把这些归结为太大的显著性。换句话说，随机序列有可能被错误地解释为一种真正的趋势（见图 7-8）。

2. 过度的锚定与过少的调整

我们在第 2 章中看到，在不确定性的情况下，人们依靠启发来简化和加速像估计概率这样复杂的认知性任务。到目前为止，我们还没有讨论过的一种启发法就是锚定效应。它是以估计量为基础的。以初步信息为基础

做出初始数量的估计称为锚定。然后, 根据额外信息对初始估计进行向上或者向下的调整。锚定效应法则看上去还是有些道理的。

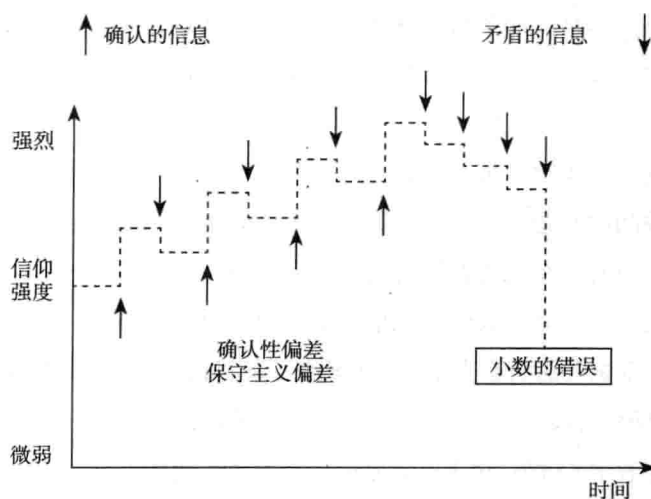


图 7-8 修正的观念

然而, 在实际当中, 当运用启发法的时候, 人们通常会犯两个错误: 第一, 初始的估计往往会受到完全不相关的锚定的强烈影响, 甚至当锚定的不相关性完全明显时, 这种情况也会出现; 第二, 以相关的锚定信息为基础做出的初始估计, 再做出后续的向上或者向下的调整会非常有限 (保守主义偏差)。换句话说, 额外信息对初始估计的影响非常小。

对利用不相关锚定的倾向可以通过一项测试人们估计密西西比河长度 (实际长度为 2 348 英里^①) 的研究中得到论证。受试者首先被问到密西西比河的长度是否大于或者小于某些通过主观臆断得出的长度。假设锚定法则设定的数字已经存在于受试者的大脑中, 当我们提出“密西西比河的长度是大于还是小于 800 英里时”, 大多数人的回答是大于 800 英里。受试者接下来被问到的是对密西西比河实际长度的估计, 其他的受试者被问到密西西比河实际的长度是大于 5 000 英里还是小于 5 000 英里。这一次, 大多数人纠正了自己的答案, 回答说密西西比河的长度小于 5 000 英里, 但是他们对密西西比河实际长度的估计是, 平均来说, 大于受试者 (被问到

① 1 英里 = 1 609.344 米。

与 800 英里做比较的受试者) 给出的估计值。这项研究显示, 受试者利用问题中已知的英里数, 然后据此进行充分调整。如果数字 800 或者 5 000 并没有影响到估计, 那么对密西西比河长度的估计就会出现相同的结果, 人们就不会去考虑这个问题是否包括与 800 英里的比较还是与 5 000⁶⁵ 英里的比较。

另外的一个实验显示, 从们会对一个明显不相关的数进行锚定。⁶⁶ 受试者这次被问的问题都不属于常识, 例如, 联合国中非洲国家所占的百分比是多少? 对每一个问题的回答包含一个百分比 (0 ~ 100%)。在每一个问题提问之前, 一个由 100 个卡槽组成的轮盘赌都会在受试者的面前转动。研究显示, 受试者的回答明显受到了轮盘赌随机得出的结果的影响。例如, 如果轮盘赌停留在数字 10 的位置上, 那么中位数的估计就是 25%; 但是当轮盘停留在数字 65 的位置上时, 则中位数的估计就是 45%。

锚定效应的启发法被认为与投资者的反应不足相关。对牛市信息的反应不足导致了资产价格仍然停留在便宜的低价位上, 而对于熊市信息的过度反应则使资产的价格收于高位。随着时间的推移, 市场暂时的低效率状态会随着价格向理性水平的转移 (趋势) 而得以解决。因此, 锚定效应有助于解释价格趋势的出现。

锚定效应可以解释前面提到的动量策略的盈利性。⁶⁷ 它以简单的技术指标为基础, 即一只股票接近 52 周高点时的情况。由于这类信息可以在报纸和各类网站上获得, 投资者可以确定 (锚定) 它。换句话说, 如果现在进行交易的股票价格接近那个水平, 那么投资者就可以固定或者锚定 52 周价格的高点。实际上, 股票就陷入了 52 周高点的怪圈当中。因此, 价格就无法恰当地反映新的牛市信息了, 股票价格就暂时性地被低估了。最终, 价格错误会随着向上的系统运动的不断提高而得到修正。

3. 对故事的锚定

投资者不仅仅会对数字进行锚定, 而且还会受到看上去引人注目的故事的蒙蔽。⁶⁸ 这两种情况的效果是相同的, 价格不能对新的信息做出有效的反应, 因此会背离理性评估。

故事之所以引人注目, 是因为“很多导致人类行动的思想是不能用数

量来表示的,但是代之以讲述故事和辩护的形式出现”。⁶⁹正如我们在第2章中指出的那样,投资者更多依靠的是有因果联系的叙述,而不是把实证进行权衡和合并来估计概率并做出决策。人们看上去需要简单的理论基础对他们的行为进行证明。⁷⁰最具有吸引力的理论基础貌似合情合理,易于理解,便于重复,因果内容充实且细节丰富多彩。这些故事对投资者是很有吸引力的,并且可以刺激买卖交易。

4. 乐观主义和过度自信

一般来说,人们对于他们自身知识的质量和精确程度都是过度自信的,因此,投资者对于他们对公共信息的私人解读具有过度自信的倾向,并且对于他们将要获得的利润过度乐观。过度自信和过度乐观的综合效应导致了投资者对他们的私人信息反应过度,而这又反过来将证券的价格推高。过度扩张的价格变化会导致价格的反转和向理性水平回归的系统性运动。

5. 小数错误(对样本容量的忽视)

小数错误是指对包含用于形成结论的样本数量缺少考虑。因此,对于判断一个数据是否是由随机的或者是非随机过程产生的,还是以小数为基础对总体参数进行估计来说,是不合理的。例如,如果在连续10次掷硬币的过程中出现了7次正面(头像),就得出出现正面的概率是0.7的结论就是无效的。而真正出现正面(头像)的概率则只能依靠从大量的掷硬币行为中推导。

当试图决定一组时间序列是否是随机的时候,忽视样本容量的投资者非常容易受到蒙蔽。例如,过快地对少量实证进行归纳,投资者也许会对公司几个季度以来产生的正利润做出公司已经建立了有效的增长趋势的错误结论。我们已经看到少数的正盈利季度集群在随机出现的盈利趋势下是很容易出现的。因此,小数错误有助于解释投资者为什么会对短期内连续出现的盈利报告做出过度反应。

小数错误可以导致两种不同的判断错误:赌徒的谬论和聚集性幻觉。这两种错误的产生取决于观察者之前对有关正在密切观察的过程所持有的观念。假设投资者看到了连续的价格变动,而这些变动实际上是相互独立

的（随机游走）。

如果投资者之前持有变动过程是随机的观念的话，那么这种错误属于赌徒的谬论范畴。在以常识为基础的情况下，我们对随机过程的触发器的显示和少量的连续性（趋势）的期望要大于对其实际发生的期望。结果，连续的正价格变动会导致天真的观察者歪曲其所做的预期趋势反转（负的价格变动）的结论。因此，在出现了5次正面（头像）之后，观察者得出了结果可能是超过50/50的结论，这种情况是真实可能的。在实际中，随机游走的序列并没有记忆，所以，强势的存在并不能改变下一个用任何方式得出的结果。我们把对反转的错误预期称为“赌徒谬误”。因为从统计学的角度来讲，当对轮盘赌受到关注的时候，天真的投资者就会犯这种错误：连续黑色的结果被认为会增加红色出现的概率这句话是完全错误的。

另外一种来源于忽视样本规模的谬误就是聚集性幻觉。当观察者之前没有对产生数据的过程是随机的还是非随机的情况形成观念时，聚集性幻觉就会出现。聚集性幻觉是对实际上属于随机游走的数据的顺序（非随机性）产生的错觉。我们再来假设有人试图通过对随机游走的过程产生的结果进行的观察来确定该过程是随机的还是非随机的（顺序的、系统的）⁷¹情形。回想一下，随机游走的小样本通常会比引起我们产生预期（篮球比赛中的火热手感）的常识显示更多的趋势（群集的）。聚集性幻觉的结果是，连续的积极的价格形态以及盈利变动会被错误地理解为合理的趋势，而这一趋势原本是随机游走中的普通走势而已。

社会因素：模仿行为、羊群效应以及信息瀑布⁷²

我们已经了解到在个人投资者的水平上，投资者行为是如何解释各种不同类型的系统性价格变动的。而这部分内容分析的是在群体性水平上，投资者行为对系统性价格变动做出的解释。与传统的投资者心理学技术分析理论相反的是，我们可以看到在某些情况下，群体行为反而是十分理性的。

当面对不确定性的选择时，人们通常会从其他人的行为中得到些许线索，并模仿他们的行为，我们把这种行为称为羊群效应（herd behavior）。

我在此需要说明的一点是，羊群效应并不是指相似的行为。当个人做出了相似的选择，而这些选择又都是独立存在的时候，很多人的相似行为就会出现。羊群效应特指源自模仿的相似行为。正如我将要指出的那样，选择模仿是完全理性的。

在源自羊群效应的相似行为和源自许多决策者独立做出同样选择的相似行为之间存在着重要的差异。羊群效应阻止了信息在整个群体中的扩散，但是并没有阻止独立决策的传播。当信息的扩散得到阻止时，投资者犯同样的错误以及价格系统性地背离理性水平的可能会增加。

为了理解羊群效应为什么会阻止信息的扩散，我们需要考虑相反的情况，即个人独立评估信息并做出自主选择。假设很多投资者正在独立地对一只股票进行评估。有些人会错误地高估股票的价格，有些人则会错误地低估股票的价格。由于这些错误都是独立的，那么二者就会相互抵消，并使平均的评估错误接近于零。现在，每一位投资者都做出了独立的评估，并且在实际中按照独特的信号操作，这就使得所有关于股票的可得信息将会得到考虑并融入到股票的价格中成为可能。即使没有单个的投资者拥有全部的事实资料，那么作为整体出现的群体，也是可以获得全部信息的。

相比之下，当投资者选择模仿他人行为，而不是依靠自己独立的决策时，所有关于股票的相关信息就不太可能在整个群体当中得到完全传播，这很有可能使得全部相关信息无法通过股票的价格反映出来。然而，单独的投资者选择模仿也许还是一种完美的理性选择。并不是每一个人都有时间或者具有专业的知识来对股票进行评估，所以像跟随沃伦·巴菲特这样一个有着良好业绩记录的投资者操作也是合情合理的。一旦巴菲特通过分析决定买入或者卖出时，投资者就会停止他们自己收集信息、评估股票的努力。这就有可能出现这样一种情况，如果巴菲特疏忽了某些因素，那么投资者这个群体也会忽略这些因素。所以，我们得出了一个结论：独立的投资者采取的行为是出于理性的，而作为一个整体的群体所做出的选择是有百害而无一利的。

1. 我们为什么要模仿

人们曾一度认为模仿的行为是顺从社会压力的结果，而早期的实证看

起来也是认可这一点的。⁷³ 在一项研究中，一个值得信赖的受试者要求对由一个假冒的受试者估计出来的线段的长度进行判断。这名假冒的受试者其实也是这项实验的一部分，他故意给出了错误的答案。尽管正确的答案非常明显，但是这位值得信赖的受试者通常情况下也是跟随着群体的答案，而不是自己给出明显正确的答案。然而，在稍后的实验中，⁷⁴ 我们把这位值得信赖的受试者与他所在的群体相隔离，而他给出的答案仍然是这个群体所持有的明显错误的答案。这个发现意味着社会压力并不能解释模仿的行为。一个更好的解释出现了，受试者看上去是依赖了社会的启发，即当一个人的判断与大多数人的判断出现矛盾时，他们就会顺从大多数人的判断。换句话说，人的行为是遵循显性规则的：因为大多数人的判断应该是不会错的。

这种行为是理性预测的结果：在每天的生活中，当一大群人在一个简单的事实上取得一致的判断时，这个群体的成员几乎是正确的。⁷⁵ 同理，我们当中的大多数人都是遵循着这个原则行事的，当一位专家告诉我们一些与常识相矛盾的事情时，我们倾向于听从这位专家的意见。所以，在充满不确定性的条件下，相信专家的建议或者追随大多数人的行为是完全合理的。

2. 信息瀑布与羊群效应

投资的不确定性使得投资者有可能会依赖于启发式信息，而不是独立决策。这对于系统价格运动的出现来说是非常重要的启示，这是因为模仿可以引发“信息瀑布”。⁷⁶ 信息瀑布是指一连串个人的或者是几个人的模仿行为。在某些情况下，模仿行为一直以来是随机选择的。

信息瀑布可以由一个假设的例子来表示。假设两家相邻的新餐厅在同一天开业。第一位来此就餐的顾客必须选择一家餐厅，那么在其做出决定之前几乎没有可以参考的信息，因此，这位顾客的选择就是随机决定。然而，当第二位顾客来就餐时，他就有了可供参考的额外信息：有一家餐厅有人在吃饭。这就使得第二位顾客到同一家餐厅就餐。第一位顾客的选择是靠猜测，而第二位顾客的选择是以不包含任何信息的信号为基础的。当第三位顾客看到有一家餐厅没有人，而另外一家餐厅有两个人时，就会促

使他做出模仿前者的选择。最后，随着独立选择的叠加，最终导致了一家兴旺一家愁的局面，而这也完全有可能导致一家提供更好食物的餐厅倒闭。

因此，最初的随机事件设定了故事朝向特定路径（餐厅1的成功）的进程，而不是向另外一条路径（餐厅2的成功）发展的方向。随着时间的推移，出现第一家餐厅火爆而第二家餐厅倒闭的可能性就大大增加了。同理，最初的随机价格运动可以引起连续的模仿投资行为，并导致长时间的大幅价格波动。

我们现在考虑一下，在许多顾客都做出了独立选择餐厅的决策情况下，将会发生什么事情（我们允许两家餐厅作为众多顾客的样本）？把他们的发现进行结合，就可以知道哪家餐厅是更好的。在信息瀑布中，模仿行为防止结果的积累，并且分享许多独立的评估结果。因此，信息瀑布限制了信息和理性选择的扩散。

信息瀑布模型显示了存在于金融市场的价格是由类似的投票（很多个人自己做出的评估）程序所决定的这一观点的缺陷。⁷⁷在信息瀑布中，投资者做出了模仿他人行为的理性选择，而不是为了获得独立的选择去花费大把的精力。投票人只做一件事，而投资者所做的事情有时却不尽相同。

3. 信息在投资者之间扩散

我们已经看到信息瀑布可以解释个人投资者的理性行为会阻碍信息扩散的原因了。这对技术分析来说尤为重要，因为信息在投资者之间的扩散速度可以解释系统性价格变动的出现。有效市场假说假定信息的瞬间扩散导致了价格几乎会在瞬间被调整这一结果。这种价格调整出现的如此之快，以至于技术分析方法无法得到利用。然而，以较慢的速度逐步扩散的信息、系统性价格变动则可适用于技术分析方法。

尽管先进的通信科技在当下流行，但是交换投资信息的首选方法还是依靠把好的故事通过以人传人的方法来实现。这种偏好是几百万年来进化的结果，这种观点已经在行为金融学家罗伯特·希勒的研究中得到了确认。他证明了即使在阅读量很大的人中间，口头传达的效果也是最强的。在聊天中出现的新的热门事件所产生的影响要比新公司高破产率的统计数字大得多。

信息在投资者之间扩散的方式通过数学模型得到了研究，而这种模型类似于那些研究疾病是如何在人群中传染的流行病学家使用的模型。遗憾的是，这些应用于投资者领域的模型的精确程度还没有达到其在生物学领域的精确程度。这种观点是通过循环想法的突变率要高于生物的变异率这一事实进行说明的。然而，这些模型以另外一种方式起到了启发作用，它们解释了故事是如何快速传播的。

一个“新”故事可以得到快速传播的原因是，它本身就是个旧故事。由于许多类似的剧本故事已经布满投资者的大脑，所以留存在投资者大脑里的那些相互矛盾的故事就显得无所谓了。例如，投资者可以一边相信股票市场是不可能预测的观点，又可以一边持有与其相反的观点。毫无疑问，大多数的投资者都曾经遇到过同时支持正反两种观点的专家。我们的大脑能够接受各种矛盾的思想，是因为眼下的这些剧本故事对于我们没有选择的必要。

然而，这件事情很快就會发生改变。不论是信息还是市场行为的轻微随机波动，都会导致深埋于投资者头脑之中的想法焕发生机，并且分散他们的注意力。

4. 投资者注意力的转变

投资者把他们的注意力都集中在了任何信息的细小变化都导致发生显著变化的给定时点上。人类大脑的组织结构在本质上具有在同一时间将注意力集中于一点，并且可以在不同的焦点之间快速切换的功能。⁷⁸ 人类的大脑已经进化到能够有选择地把精力从关注流入世界的信息洪流转移到其中的细小支流的程度。这一过滤的过程是人类智慧的标志，以及在得知结果之前，我们根本不知道哪个细节是值得我们关注的。

大脑用来过滤相关信息时所使用的无意识法则就是从观察他人行为中得到线索的法则。换句话说，我们假设能够吸引其他人注意的事物必须是那些也吸引我们注意的事物，根据经济学家罗伯特·希勒所说的：“引起社会关注的现象是行为进化的伟大产物，这对于人类社会的运转来说是非常重要的。”⁷⁹ 尽管得到公众的关注具有重要的社会价值，但是因为它促进了

合作行为，所以是负面的。它会导致整个群体持有不正确的观点，并采取类似的错误行为。

在相反的情形下，投资者独立地形成了他们自己的观点，而错误的关注则是随机且可以自动抵消的。尽管这种思维模式很难围绕着普遍的价值观和目标形成共同努力的局面，但是它也减少了忽视重要细节的可能性。充满剧烈波动的市场运动可以轻松抓住投资者群体的注意力，并且鼓励他们以扩大运动变化的形式来操作，即使这种操作根本就没有原因。

罗伯特·希勒研究了抓住投资者注意力的充满剧烈价格运动以及引发羊群效应的方式。即使是那些以有效的统计学特征为基础，用系统的方法挑选股票的机构投资者，也会具有仅仅因为股票经历了快速的价格上涨而买入股票的倾向。⁸⁰ 这种操作机制可以说明被吸引者及逼真的初始价格变动如何引发市场繁荣这一现象。此外，希勒的研究显示了投资者通常不会关注导致他们进行操作的剧烈的价格变动。

5. 反馈在系统价格变动中的作用

投资者之间存在的社会互动产生了反馈。反馈是指系统输出以输入的形式重新进入系统的渠道（见图 7-9）。图 7-9 中对两套系统进行了比较，即没有反馈的系统和有反馈的系统。

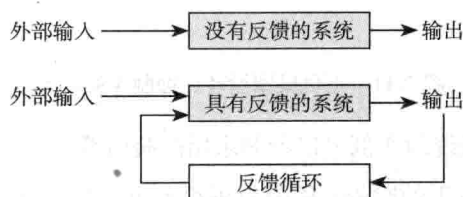


图 7-9 反馈：输出转变为输入

反馈有正反馈和负反馈两种类型。它们在思维方式上存在差异，如下：在负反馈的情况下，用系统输出乘以负数就得出系统输入的反馈；而在正反馈的情况下，乘数则为正数（见图 7-10）。

负反馈具有抑制系统行为的效果，并推动输出回归到均衡水平。正反馈则具有相反的效果，它会扩大系统性行为，并将输出推向远离均衡的水平。这种情况如图 7-11 所示。

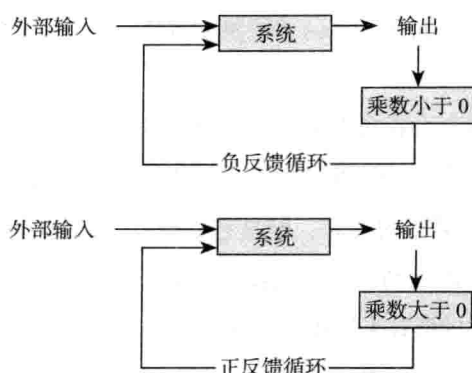


图 7-10 正反馈和负反馈循环

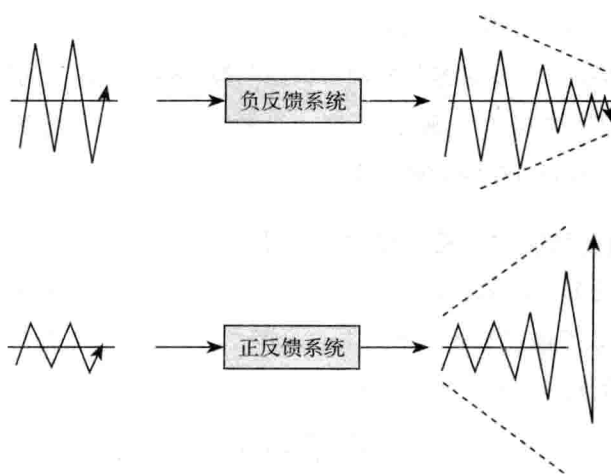


图 7-11 正负反馈循环：抑制 VS. 扩大

关于负反馈系统的实例可以参照家用的暖风系统。当温度降低到期望值或是均衡水平以下的时候，恒温器就会打开暖气，让系统的输出返回，也就是屋子的温度回归理想水平；当温度高于期望的水平时，空调就会打开，把温度再降下去。这种模式就需要将系统输出以输入的方式再输入回去。这种方法可以让系统通过自动调节的方式将温度达到均衡水平。

正反馈的例子可以用人们靠近扩音器时使得公共扩音系统不断上升的音量为例。扬声器（输出）发出的正常的嗡嗡声通过扩音器返回到音频系统，然后再由音频系统将其扩大。随着每一次的循环，声音越来越大，最终成为刺耳的声音。正反馈也被称为雪球效应（snowball effect）或者恶性

循环 (vicious cycle)。

金融市场，与其他的组织自我调控系统一样，依靠的是负反馈与正反馈之间的均衡。⁸¹ 套利导致了负反馈。价格过高或者价格过低引发了套利交易，并且使价格向理性水平回归。当投资者依靠模仿行为而不是独立做出决定时，正反馈就会发生。在这种情况下，投资者就会跳过初始的价格变动，在第一个强势信号出现以后买入，或者在第一个弱势信号出现以后卖出，最后，把最初比较微弱的价格变动发展成为大规模的趋势。有一种被称为“趋势追踪法”的技术分析方法，它的盈利就是依靠大规模的价格变动，并且在正反馈成为主流的这段时间内会产生明显的效果。

正反馈可以导致价格上升或者价格下跌的恶性循环，即使由整个事件引发的最初价格变动被证明是合理的，也会使价格远远超过其理性水平。换句话说，由于正反馈的放大效应，价格趋势可以远高于或者低于理性水平。

正反馈也可以从人们接受他们对最近变动的预期中得到。尽管当环境改变时，人们对预期做出调整是合情合理的，但是人们并没有完全这样去做。有时候，他们并没有进行充分调整（保守主义偏差），而在某些时候，他们又对环境的变化十分敏感，并且他们对预期做出的改变比之前经过调整的预期还要大。投资者特别容易对那些引人注目的新信息做出过度反应。不论是单一的大幅价格变动，还是连续的类似价格变动，都可以引起投资者对预期做出较大的改变。这些价格变动可以满足自身的要求，因此可以导致价格出现泡沫或者泡沫破裂的趋势。

有时候出现正反馈的判断偏差是出于心理账户的原因。这种对财富非理性的思维倾向就好像它属于不同的账户一样，需要分别对待。以前通过投机冒险获得的财富可以归于即时行为账户，而对待通过资产净值账户的叠加获得的财富则添加了保守主义色彩。理性地说，投资者的全部财富，不论他是怎么获得的，或者说是通过投资哪种资产而获得的，都应当以同样的方式对待，1 美元就是 1 美元。然而，从最近上涨的市场中通过投机获得的财富，通常都要为下一次的投机活动预留，这样就可以为市场的上涨添加燃料。

在某些情况下，不断增加的预期也会像添加了燃料的泡沫一样破裂。常识告诉我们，某些分析师也在强调，当市场出现崩盘的时候（负泡沫），价格泡沫也会随之终结。有些时候是这样的，但是出现其他泡沫终止形式也是有可能的。把反馈的循环合并之后运行的电脑系统模拟显示，泡沫不仅在很多中间停顿的状态下以锯齿状的形态发展，而且还以相似的方式收缩。如果 20 世纪 90 年代末出现了价格泡沫，那么它就应当以很多中间停顿的方式而结束。

自行组织的庞氏骗局

价格泡沫的反馈理论之所以如此的吸引人，是因为它表面上看似合理，再加上它所列举的历史实例与其自身的正确性相符合。然而，事实很难证明一个仅包括加强投资者焦点、模拟行为以及夸大了的投资者信心在内的反馈原理可以在金融市场中真正得到应用。

然而，耶鲁大学的经济学家罗伯特·希勒认为，庞氏骗局的金字塔式欺诈这一具体案例，为确认反馈理论提供了实证。这场骗局，是以臭名昭著的查尔斯·庞氏（Charles Ponzi）而得名的。庞氏在 20 世纪 20 年代虚构了一种想法，即他承诺通过某些商业冒险或投资支付给投资者高额的收益率。然而，投资者的资金并没有像他所承诺的那样投入使用，而是通过向新的投资者的承诺来支付之前那些投资者的收益。庞氏这么做出于两个目的：让他的投资看上去合法化，以及让第一批的投资者把自己成功的故事传播出去。当这些故事广为人知的时候，就刺激了新的投资者前来加入庞氏。而他们所投入的资金反过来用于支付前面那些投资者的收益，而那些投资者现在仍然四处奔走相告他们的成功。这种用彼得的钱来支付保罗的收益的循环成功地证明了公司的持续性和增长性。受骗者的数量沿着病毒一样的曲线在不断增加。最终，这笔投资达到了饱和状态。没有受到感染的投机者对该笔投资的资金供应消耗殆尽，已经不再有新的投资来为最后一批投资者的投资支付收益了。此时，骗局失败了。

近来，一位来自阿拉斯加某小镇的家庭主妇在 6 年的时间里，通过承诺支付 50% 的收益率，获得了大约 1 000 万～1 500 万美元的收入。她的

行为据说是以大航空公司叠加累计的未使用的航空里程为基础的。实际上，这种行为与庞氏骗局非常相似，只不过庞氏骗局是以国际邮票优惠券为基础的。

庞氏骗局又进化成为了一个典型形态。起初，投资者对此表示怀疑，并且只投入了少量资金，但是一旦第一批投资者收获了利润，并把他们成功的故事告诉其他人，一大批投资者就会因此而获得自信，并且加快资金注入势头。但是对此持有理性怀疑态度的人频繁地发出警告，但是人们贪婪的本性却把这些话当作耳旁风。然而，那些理由充分的带有批评性的争论却无法与赚得盆满钵满的账户一争高下，最终，飞蛾扑火的事情发生了。

希勒主张，在不需要发起人具有欺骗性的阴谋下，⁸² 投机性的泡沫自然会产生庞氏骗局，而庞氏骗局又会自然而然地出现在金融市场当中。诸如自行组织这种现象在复杂的系统中是非常常见的。我们之所以不列举虚假的故事，是因为股票市场中总会出现一些流言蜚语。价格上升的趋势曾经在庞氏骗局中呼唤着那些获得利润的早期投资者。这些信息随后被华尔街的销售员进行了扩展，因为他们再也不需要编造谎言了。他们只重点描述上涨的潜力，却避而不谈下跌的风险。庞氏骗局与投机性泡沫的平衡点在希勒的观点中已经很明显，也就是说，举证的责任落在了那些拒绝承认相似性的人的身上。

相互矛盾的行为金融学假说

一个数学模型可以把关于世界某些方面的科学假说进行量化，所以这就使得对未来的观察结果进行定量预测成为了可能。当然，这些被提出的模型希望它的预测可以与上述的观察结果相一致。然而，如果二者之间出现了不一致，则真正的科学家会时刻准备着制定一个新的模型，而这个模型既可以解释不一致的观察结果，又可以做出能够被未来的观察结果进行测试的额外预测。由于我们的知识永远都不能是完整的，所以这一精确的循环永远也不会停止。

伪科学看上去很像合情合理的科学假说，但是这些都只是表面上的。科学的假设虽然是简洁的，但是它对范围广泛的观察结果进行了解释说明。

伪科学则倾向于复杂化。科学假说做出的是简明的预测，而且是可以被未来的观察结果所检验的。而伪科学报告做出的则是模糊的预测，虽然它与之前的观察结果相一致，但是永远也不会做出能够被明确反驳的具体预测。因此，不论它距离标准有多远，伪科学报告只能存活在那些容易上当受骗的人的大脑中。

在此基础之上，行为金融学具备了科学条件，即使这一理论尚不成熟。虽说对单一的理论并没有形成一致的看法，但是已经取得了很大的进展。在过去的10年间，行为金融学的实践者们已经提出了很多假说来解释投资者行为与市场动态的主要特征。因为技术分析实际上是没有理论基础的，所以他们对系统性价格变动的出现提供了理论基础，这对于技术分析而言是非常重要的。

行为金融学假说已经形成了数学模型，所以它们在假说正确的条件下，可以产生测试市场行为的预测。如果市场没有显示经过预测的行为，就等于向反驳打开了大门。例如，如果模型预测趋势反转的出现类似于那些实际观察的市场数据，那么该假说就会得到确认；然而，如果模型所预测的局面与实际的市场行为相矛盾，那么这个假说就会被反驳。这样就使得行为金融学的假说具有了科学意义。

回顾一下，有效市场假说的随机游走模型做出的系统性价格运动是不可能市场走势当中被观察的这一预测。相反，行为金融学假说后来预测了系统性地价格运动的出现，并且对其出现的原因进行了解释。即使没有人假设提出市场行为的综合理论，他们也会对技术分析提出重要的中心思想：我们有理由相信金融市场并不完全是随机游走的。

1. 对公共信息做出的偏差解释：巴伯里斯 (Barberis)、施莱弗 (Shleifer) 以及维什尼 (BSV) 假说

实证研究显示投资者对新闻报告表现出了两种明显的错误：有时他们对新闻反应不足，结果导致了价格的变化要比新的基本面应该发生的变化小得多；另一些时候，投资者又会对新闻反应过度，并导致价格的变动要大于经过调整后的价格变动。然而，最终这些定价失误会通过系统性的价格变动向能够准确反映新闻信息的价格方向运动。

但问题是,为什么投资者有时会对新闻做出过度反应,或者是反应不足呢?一项由巴伯里斯、施莱弗以及维什尼(BSV)⁸³提出的假说对此进行了说明。首先,BSV认为这里面包含两个认知性错误:保守主义偏差和对样本容量的忽视。其次,他们主张在特定时点进行的错误操作主要取决于环境因素。

现在简单回顾一下,保守主义偏差描述的是,以前的观点并没有因为新信息的出现而发生太大的改变,现存的观点是不会改变的。对样本规模的忽视也被称为“小数错误”,这种形式的错误具有从少量的观察结果当中得出重要结论的倾向。因此,保守主义偏差是小数错误的对立面。前者对新实证的重视程度不够,而后者又对新实证反应过度。

当投资者对新闻反应不足时,保守主义偏差便开始发挥作用。当出现利好消息时,保守主义偏差会把价格压得很低;而在面对利空消息时,它把价格抬得很高。然而,当投资者逐渐意识到新信息的真实意义时,价格便会系统性地逐步走向能够恰当对信息做出反应的价格水平。因此,有些价格趋势是可以用投资者对信息的反应不足来解释的。

小数错误解释的是投资者对信息做出的过度反应。例如,如果投资者观察少量的正收益率变化,他们就会草率地得出有效的上涨形态已经确立的结论。这一结论就会过度地刺激投资者马上买入证券,并把其价格推高到在正常价格水平以上的价位。投资者有时也会对一些负面的观察结果反应过度,并且把价格压低。

由BSV提出的数学模型做出了下列假设:①市场上只有一只股票;②该股票的收益率盈利流是以随机游走为基础的;③投资者并没有意识到这些;⁸⁴④股票的盈利流在增长趋势阶段和均值回归阶段(震荡)之间反复,投资者就是在这种错觉下进行操作的,但是他们并没有注意到盈利也可以是以随机游走为基础的;⑤因此,投资者相信他们的工作就是在任一时点上及时判断出到底哪种形态是有效的。

让我们看看由BSV所虚构的投资者是如何对新盈利信息的发布做出回应的。在给定的任何时间点,投资者都持有一种信念,即盈利是出现在上涨趋势和均值回归过程这两种非随机状态之中的。这里并没有把随机游走

的可能性考虑进来。回顾一下，随机游走所显示的均值回归要小于真正的均值回归过程，而且其持续的时间也要比真实的趋势过程短。我们再来看这个例子，假设新闻的发布与投资者持有的观点相矛盾。例如，如果投资者坚信股票的盈利来自上涨趋势，但是却出现了负面的盈利报告，这显然是出乎了投资者的预料。同理，如果投资者认为盈利来自均值回归状态，那么当出现类似的连续变化时，投资者也会感到意外。根据 BSV 假设，这两种出人意料的情况都会由于保守主义偏差而导致投资者的反应不足。换句话说，不论投资者之前所持有的观点是上面所述的哪一种，与其观点相矛盾的新闻都不会被重视，因此就会导致价格对于新信息的反应不足。然而，如果与之前持有的观点不一致的现象继续的话，投资者就会相信之前的状态已经改变了，不论现在的状态是从均值回归到上涨趋势，还是从上涨趋势到均值回归。在这种情况下，BSV 假设预测投资者将会从坚持以前的观点和反应不足（保守主义偏差）的错误当中转变到承认小数错误和反应过度的错误中来。因此，BSV 模型对两种投资者所犯的错误，以及纠正这两种错误形式的系统性价格变动做出了解释。

我把由 BSV 模型对行为类型所做的预测绘制在图 7-12 中。我在此处假设投资者最初相信盈利来自均值回归过程，因此，他们期望负面的报道会在正面的盈利报告之后出现。出于这种原因，他们对前两个正面的盈利报告反应不足。然而，在第三个正面的盈利报告出现之后，投资者拒绝其之前持有的观点，而接受盈利来自已经确立了的上涨趋势的观点，这一点在随机过程（三个头部连成一行）中并不罕见。所以，投资者现在期望未来出现的盈利报告是正面的。这种情况就导致了投资者的过度反应，从而把股票的价格抬到理性水平之上。由两种投资者错误导致的系统性价格变动如图 7-12 所示的灰色部分。

BSV 模型也显示了反馈循环是如何从小数错误中出现的。当投资者看到盈利连续多次进入相同的方向时，他们就会从对趋势的质疑转变为假设它是真实的状态。⁸⁵ 然而，保守主义偏差最初阻碍了对新观点的接受。小数错误和保守主义偏差之间的相互作用决定了投资性反馈的发展程度。⁸⁶ 对 BSV 模型的模拟显示了多种实际的市场行为，如对意外收益进行趋势

追踪、价格惯性、长期反转以及对市净率等基本面比率的预测力。因此，BSV 模型解释了投资者对公共信息的偏差解读是如何对存在于实际市场数据中的现象做出预测的。这就要求模型抓住某些金融市场的实际动态。

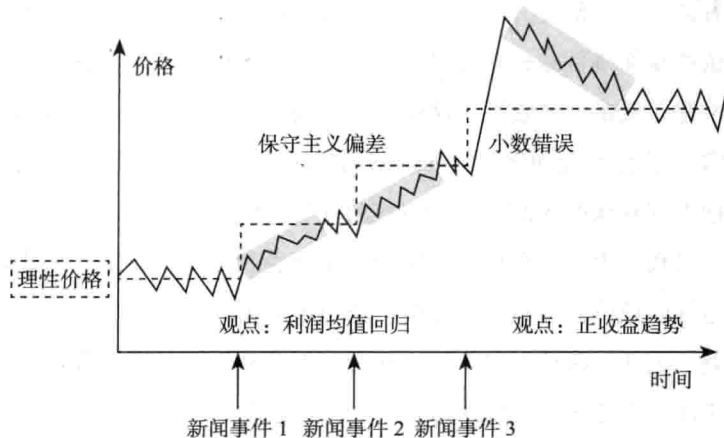


图 7-12 BSV: 伴随着过度反应而来的反应不足

2. 私人信息的偏差解读: 丹尼尔 (Daniel)、赫舒拉发 (Hirshleifer) 以及苏布拉马尼亚 (Subrahmanyam) (DHS) 假说

由丹尼尔、赫舒拉发以及苏布拉马尼亚 (1998, 2001)⁸⁷ 三人提出的另外一种假说对诸如趋势反转和趋势持续 (势头) 等系统性价格变动给出了不同的解释。与 BSV 不同的是, DHS 假说是以投资者对私人研究的偏差解读为基础的。此外, DHS 强调了一些不同的认知性错误: 禀赋偏差、确认性偏差 (参见第 2 章) 以及自我归因偏差 (参见第 2 章)。

私人研究是指投资者对公共信息进行独立的推断解读 (专有研究)。由 DHS 提出的这一假说预测投资者将会对自己所做的私人研究过度自信, 并且会对该研究具有反应过度的倾向。这种预测是以禀赋偏差为基础的, 即人类有把我们自身或者创造的价值抬到很高位置的天赋。这种反应过度的表现把价格推高到远离理性价格水平, 最终导致价格出现反转, 并向着理性价格回归。所以, 在 DHS 看来, 价格反转是投资者对自己独立推断信号过度自信的结果。

DHS 还对由确认性偏差和自我归因偏差这两种认知性偏差引起的价格

惯性（非随机价格趋势）做出了解释。他们主张，由于新闻与价格变动以一种偏差的方式影响着投资者，才使得价格惯性出现了变化。由于确认性偏差的缘故，那些能够确认投资者以前所有观点的新闻或者价格变动（即以投资者通过自己研究并发现信号为基础的观点），就会得到格外的重视。得到确认的新闻事件或价格变动具有增加投资者信心（认为他们的私人研究实际上是正确的）的效果。这就引发了投资者采取进一步的行动（买入或者卖出），这种行为加强了当前价格趋势的动力。自我归因偏差也会导致投资者把得到确认的新闻或者价格变动视为他们私人研究成功的标志。回顾一下，自我归因偏差会使人们因为结果而受到好评，即使这些结果的出现包含运气的成分，但是人们也会出于拒绝承担坏的结果而把原因归结为运气太差。因此，确认性偏差和自我归因偏差都会引起投资者对已经得到确认的实证反应过度，从而放大了市场的正反馈。

相反，如果公共信息或者价格行为与投资者的私人研究相矛盾，那么投资者就不会对此非常重视（确认性偏差），并且投资者还会把这一切归结为运气太差，而不是归结为其私人研究（自我归因偏差）的不足。因此，不论新闻或者价格变动与投资者的私人研究相一致还是相矛盾，DHS 模型都会做出投资者对于他们私人研究的自信不断增加的预测。这种对新信息独特的处理意味着最初的过度自信通常会带来更大的过度自信，因此，就会产生股价惯性。⁸⁸ 然而，在出现了连续的与公共新闻相矛盾的信号时，投资者对他们的私人研究的信心将会得到削弱（小数错误），价格趋势也进而出现反转。

DHS 和 BSV 都预测投资者有时会过度乐观，有时会过度悲观。此外，当一连串的反面信息出现的时候，投资者的这些观点就会倾向于反转。如果 DHS 和 BSV 是有效的，那么通过利用趋势反转制定的策略赚取的利润就应当集中在信息发布时的那个时间点上。这一预测结果已经得到了确认，并且让 DHS 和 BSV 两种假说都面临着被反驳的境地。⁸⁹ 实际上，大部分通过买入价值型股票（低市净率、低市盈率等）和在过去的 3 ~ 5 年中表现弱势的股票所赚取的超额利润都是发生在靠近盈利性信息发布的时间点的。

3. 新闻交易者和趋势交易者：HS 假说

我们曾经描述过，某些价格惯性的案例可以通过捕捉对新闻初始的反应不足的基本面价格来解释。另外一种趋势可以用正反馈来解释，即价格的上涨刺激了加码买入，而价格的走弱则刺激了加码卖出。⁹⁰ 因此，有些投资者通过最近的价格变动，而不是通过新闻来寻找信号。

站在技术分析的角度上，由正反馈引起的趋势可以被视为减少对引发趋势追踪信号敏感性的逐步的串联。最初的价格变动（受到启发的或者是随机的新闻）触发了最为敏感的趋势信号（非常短的移动平均线交叉信号），由这个信号所触发的交易足以把价格推动到某种不敏感的趋势信号水平（中期移动平均线），而它反过来又可以引发更加不敏感的信号。无论用哪种原理解释，投资者的行为都是相互关联的，并由此引发羊群效应。

由洪（Hong）和斯泰因（Stein）⁹¹ 提出的假说解释了 3 种市场现象：惯性、反应不足以及反应过度。为方便工作，他们假设把投资者分为两种，即新闻交易者（基本面分析派）和趋势交易者（技术分析派）。HS 假说主张，在这两种投资者之间的交易产生了正反馈，也就是价格惯性。HS 假设不止一种类型的投资者是以他们知识的局限性这种观点作为交易基础的，这必然会把他们的注意力限制在可得信息受到限制的部分。新闻交易者把他们的注意力都集中在基本面的研究上。根据这些信息，他们推导出对未来收益率的私人判断。反过来，趋势交易者把注意力都集中在了过去的价格变动上，并从中推导出对未来趋势的私人预测。

洪和斯泰因假设新闻交易者很少，或者根本不关注价格变动。结果，他们对新闻的私人评估便不能够快速地在整个新闻观望者的群体中进行传播。为了更好地理解对基本面的特别关注是如何阻碍信息的扩散的，我们来考虑一下，如果新闻观望者对价格行为予以更多的关注时会发生什么情况。对价格变动的观察会让他们对其他新闻观望者持有的私人信息做出某种程度的推断。例如，价格上涨也许暗示着其他新闻观望者把这些新闻视为利好，而价格的下跌则意味着他们把新闻解读为利空。正如我们已经看到的那样，任何落后于包含着新信息的价格的投资者行为都会导致价格趋势。也就是说，在没有对新信息做出及时的价格调整的情况下，价格会

逐渐地向经过新信息证明过的更高或者更低的价格水平移动。为了让惯性（趋势追踪）策略发挥作用，价格必须以一种足够渐进的方式进行改变，从而让以初始价格变动为基础的信号对另外的价格变动做出预测。因此，新闻交易者无法从价格行为中提取任何信息来减缓价格对新信息的调整速度。

现在，我们来考虑由趋势交易者（趋势追踪者）导致的市场动力造成的影响。HS 假说认为趋势交易者只关注价格行为产生的信号。趋势追踪则假设新出现的价格趋势是重要的新基本面信息开始在整个市场中扩散的标志。趋势交易者确信这种扩散是不完全的，且价格趋势在初始的趋势信号出现之后将会得到继续。趋势追踪者随后采取的行动使得价格持续运动，因此就产生了正价格反馈。最终，随着更多趋势追随者的加入，价格就会继续偏离调整后的基本面。由于趋势交易者无法判断新闻的扩散程度以及投资者对此的理解程度，对冲的情况就会出现。当新闻得到了全部的扩散，并且人们都根据它进行操作的时候，就不会再有其他价格变动了。因此，HS 认为，来自趋势交易者的正反馈，产生了理性价格水平的过冲现象，这就为价格向以基本面为基础的理性水平的趋势反转创造了条件。

4. 本质内容

行为金融学为我们提供了很多可以进行检验的假说，这些假说试图对系统性地价格运动进行解释。如果真是这样的话，那么到目前为止还没有哪种假说被证明是正确的。毫无疑问，还会有新的假说提出。只要能够产生可以检验（可以进行证伪的）的预测，他们就很满足了。对技术分析来说最重要的是，这些理论提出的系统性地价格变动也应该出现在金融市场。

发生在有效市场条件下的非随机价格运动

上面的内容解释了系统性地价格变动是如何出现在非完全有效市场中的，然而，即使市场是完全有效的，对技术分析来说，系统性地价格变动也不会就此消失。有一种情况对存在于完全有效市场中的系统性地价格运动来说是最理想的。

系统性价格运动与市场有效性共存

有些经济学家主张，即使市场是完全有效的，系统性价格变动也是可能存在的。⁹²换句话说，非随机价格行为与价格预测的可能性的存在并不意味着金融市场就是无效的。⁹³

其他经济学家则陈述了更加强有力的案例，他们主张可预测性不仅在某种程度上是可能存在的，而且对于市场的正常运转也是非常重要的。在《华尔街的非随机漫步》（*A Non-Random Walk Down Wall Street*）一书中，罗闻全和克雷格·麦金利认为，“可预测性是资本主义齿轮的润滑剂”。⁹⁴而格罗斯曼和斯蒂格利茨也提出了类似的观点，⁹⁵他们主张有效市场必须通过提供获利机会来刺激投资者加入到昂贵的信息处理和交易活动中来。正是由于这种特殊的活动才使得价格靠近理性估值。

实际上，只有在一种特殊的情况下，市场的有效性才会意味着价格的变动必须是随机的，也就是不可预测的。当所有投资者对风险具有相同的态度时，这种情况就会发生，而这种情况并不是表面上看似合理。现在有没有一种能够让所有的投资者具有相同态度的风险类型呢？为了理解这种假设是多么的牵强，想象一个让所有人在对待风险时都能够具有同样态度的世界。我想这个世界要么是存在很多的试飞员和炼钢工人，要么就是从事这类工作的人存在严重不足；要么这个世界中的每一个人都投身于跳伞以及与短吻鳄摔跤的运动中，要么就是没有人愿意这么干。

更加现实的假设是一个由面对风险存在各种不同态度的投资者所组成的世界。在这样的一个世界中，有些投资者要比其他人更厌恶风险，他们还愿意让其他投资者接受额外风险。同理，那些能够容忍风险的投资者会寻求机会并为此承担额外风险。在这样的世界中所需要的是一种可以将风险从风险厌恶向风险容忍转变，并且对那些愿意接受上述风险的人提供补偿的机制。

让我们现在先远离一下市场，来考虑风险厌恶对喜好风险进行补偿的机制——保险费。房主担心发生火灾，所以他们就进入了与火灾保险公司的交易之中。房主把火灾产生的金融风险转嫁给了保险公司，而保险公司反过来收取保费。如果保险公司能够预测到房屋失火的概率（大量的房屋，

即大数法则)，而收取的保费又可以补偿火灾所带来的风险并产生一小部分盈利的話，这笔生意就是有利可图的。从本质上讲：保险公司提供了风险接受服务，并且为此做出赔偿。

金融市场的参与者已经暴露出了各种各样的风险，而他们在容忍风险方面却各不相同。这些因素刺激了将风险从一方转嫁给另一方的交易。在这种情况下，投资者愿意忍受更大的风险，并提出了以更高的收益率作为补偿的合理要求，而不是提出无风险投资的要求，就像试飞员和炼钢工人为了他们承担的额外风险，可以要求丰厚的报酬那样。在金融市场中，我们把对接受不断增加的风险所给予的补偿称为风险溢价（risk premium）或者经济租金（economic rent）。

金融市场提供了各种类型的风险溢价。

- **股票市场风险溢价：**投资者为新企业的成立提供营运资本，并且通过获得高于无风险利率的收益作为导致企业破产、经济衰退等风险的补偿。
- **商品和货币套期保值风险转移溢价：**投机者假设期货的多头和空头头寸可以让商业性套期保值者（商品的使用者和生产者）拥有摆脱价格变动风险的能力。而对他们的补偿则是以能够被简单的趋势追踪策略所捕捉到的盈利性趋势的形式实现的。
- **流动性溢价：**投资者选择对流动性较差的证券以及被那些短期需要现金的投资者所超卖的证券进行空仓处理。至于补偿，则是通过从持有的非流动性证券上获取更高的收益，或者是通过买入最近比较弱势（从事反潮流策略）的股票而获得短期收益等方式获得的。
- **为提升市场有效性的信息或者价格发现溢价：**这种溢价要对做出买入/卖出决策的投资者进行补偿，因为他们的决定使得价格向理性水平回归。由于这些决策是建立在复杂模型的基础上的，我们还可以把这种溢价称为“复杂性溢价”。⁹⁶例如，某种策略采取卖出被高估的股票，买入被低估的股票的策略来帮助价格向理性的价值（价格发现）回归。

有关收益率的利好消息可以以风险溢价的方式持续到未来。当然，发现市场的有效性是每一个分析师的梦想，因为通过市场的有效性获得的收益不会承担额外风险。⁹⁷ 市场的无效性只是暂时的，它迟早会被勤劳的研究者发现，并且彻底消失。然而，被认为是合理的风险溢价的收益率则很可能以服务报酬的方式持续下去。

那么风险溢价是如何解释系统性价格变动的呢？这种价格变动可以为应该得到赔偿的那些接受风险的投资者提供工作原理。那些愿意为想强烈卖出股票的股票持有者提供流动性的投资者最终会在短期内把股价提升上去。这些股票的表现与反潮流策略⁹⁸的表现是相同的。

因此，在有效市场的情况下，通过技术分析策略获得的盈利也可以被理解成为风险溢价。总的来说，对获利所做的补偿会影响其他投资者或者市场。换言之，技术分析信号之所以可以产生利润，是因为那些投资者抓住机会承担了其他投资者希望摆脱的风险。

在此，我需要指出的一个事实是，愿意忍受风险的投资者并不能保证他可以获得盈利，风险承担者必须明白这一点。正如我们之前提到的炼钢工人和试飞员有可能不会活到那么大岁数来领取包括风险任务溢价在内的工资那样，一个草率的风险溢价寻求者也不会在这个游戏里待太长时间。

对冲风险溢价与商品期货收益

这部分内容讲述的是，在商品市场中的趋势追踪者获得的利润是如何用转移风险溢价得到解释的。商品期货市场所承担的经济职能在本质上是区别于股票与债券市场的。股票与债券市场为公司提供了一种机制去获取股权和进行债务融资，⁹⁹ 并且为投资者提供了一个投资其资本的渠道。由于股票和公司债券投资把投资者暴露在超过无风险利率（政府短期无息国库券）的风险之下，那么就要用风险溢价来补偿投资者，即股权风险溢价与公司债券风险溢价。

期货市场的经济学职能与融资无关，但是却与价格风险有关。价格的变动，特别是大幅的价格变动，是生产或者使用商品的企业存在的风险与不确定性的一个来源。期货市场为这些企业提供了一个被称为商业套期保

值者的方法，它可以把价格风险转嫁给投资者（投机者）。

表面上我们仍然对商业套期保值者需要投资者来承担他们的价格风险感到迷惑。因为有些套期保值者，比如说种植小麦的农民需要将小麦卖出，而像需要使用小麦做面包的公司就需要买入小麦，那为什么套期保值者之间会相互矛盾呢？他们这么做，是因为他们的对冲需求通常是不平衡的。有时候生产小麦的农民出售的小麦超过了面包公司和其他需要购买小麦的商业使用者的需求，而在另外的一些时候，这种情况却恰恰相反。因此，在农民的供给与面包公司的需要之间常常会出现缺口，这就是引起第三方需要买卖小麦的市场参与者加入进来的原因。也就是说，他们愿意填补这个由生产商和用户造成的供需缺口。

供需缺口的存在可以通过供需法则预测：随着小麦价格的上涨，供给上升而需求下降；而当价格下跌时，供给下降而需求上涨。美国商品期货交易委员会（Commodity Futures Trading Commission, CFTC），对商业套期保值者的买卖情况进行追踪，并确认这种情况是否确实出现在期货市场。在价格上涨期间，商业卖出保值期货要大于买入保值期货（套期保值者是净卖家）；在价格下跌期间，商业买入保值期货要大于卖出保值期货（套期保值者是净买家）。

在期货市场当中，总卖出量必须与总买入量相等，就像在房地产市场中出售房屋的数量必须与购买房屋的数量相等一样，不可能出现第三种情况。考虑到价格的上涨趋势刺激了商业，卖出保值期货要大于买入保值期货，这时就需要额外的买家在市场上涨时进行调节，从而达到卖出保值期货占据优势的目的；而在价格下跌趋势当中，大量的买入保值期货引发了对额外卖方的需求。

进入价格趋势的投机者，愿意在价格上涨时通过买入的方式来面对卖出套期保值的空缺需要，并且愿意在价格下跌时进行卖空交易，从而面对买入套期保值者的需要。换句话说，期货市场需要投机者扮演趋势追踪者的角色。正是市场中这双看不见的手创造了一种趋势来为投资者填补这些缺口的行为进行补偿。

然而，趋势追踪者却暴露了显著的风险，如大多数趋势追踪系统在这

一时期所产生的不具盈利性的信号在 50% ~ 65%。当市场没有趋势的时候,这种情况就会出现,这在大多数时间里都是存在的。在这些时间里,趋势追踪者经历了股票的下跌。在股票下跌期间,作为一个成功的趋势追踪者来说,造成交易资金超过 30% 跌幅的情况是非常罕见的。因此他们需要强烈的动机去忍受下跌的风险。当出现这种情况的时候,这种动机虽然不能保证能够从价格趋势中获利,但也是一种可以尝试的机会。换句话说,大规模的趋势运动为那些资本充足的和能够正确利用杠杆作用的趋势追踪者提供了盈利机会。

趋势追踪收益率的 Mt.Lucas 管理指数

随着 Mt.Lucas 管理指数 (MLM) 的建立,由商品趋势追踪者引起的风险溢价得到了量化。由一个非常简单的趋势追踪公式获得的收益率保持着一项历史纪录,换句话说,它假设这是一个不具备专业知识的复杂的预测模型。如果真是这样的话,那么它就没有理由被称为基准指数,因为基准指数的作用就是对投资的资产类别并不具备特殊的技能而产生的收益率进行预测。

由 MLM 指数获得的经过风险调整的收益率提出,商品期货市场包含着可以用相对简单的技术分析方法来利用的系统性的价格变动。MLM 月度收益率指数可以追溯到 20 世纪 60 年代,当时它可以把 12 个月的移动平均交叉策略应用于 25 个商品市场。¹⁰⁰ 在每个月的月末,将对应于每一个商品市场就近的期货合约价格与其 12 个月的移动平均值进行比较:如果价格大于移动平均值的话,那么就对下个月的市场做多头头寸处理;如果价格低于移动平均值的话,则持有空头头寸。

MLM 指数的年化收益率和风险,¹⁰¹ 是通过年利润的标准差来衡量的,它与其他几种资产类别的年化收益率和风险的比较如图 7-13 所示。

经过风险调整的超额收益率或者夏普比率与 MLM 指数,还有其他资产类别的比较如图 7-14 所示。趋势追踪者所赚取的风险溢价在某种程度上要好于其他的资产类别基准。¹⁰² MLM 指数对期货市场以系统性价格变动的形式提供的补偿提供了一些实证。

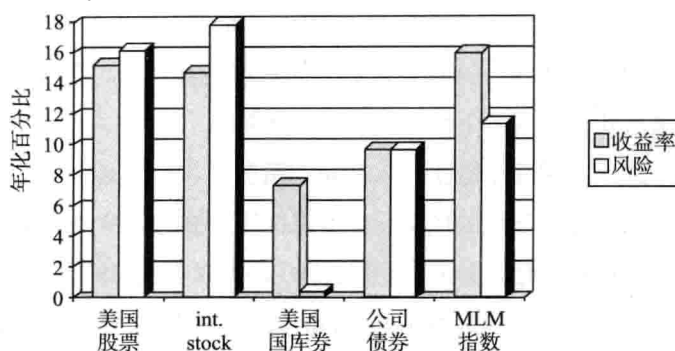


图 7-13 5 种基准的收益率和标准差

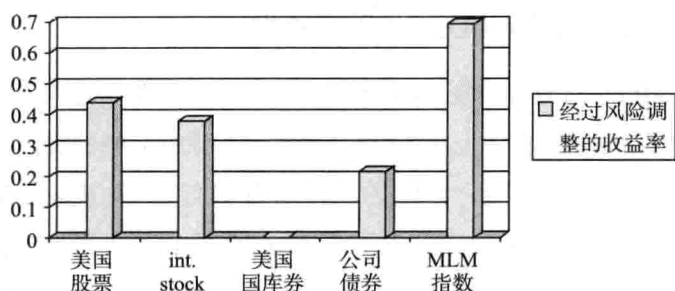


图 7-14 经过风险调整后的超额利润

如果风险转移溢价可以有效地解释趋势追踪者的收益率，那么对股票的趋势追踪就不能和商品期货市场的趋势追踪相提并论了。股票市场的趋势追踪者并不提供风险转移服务，有数据显示股票市场的趋势追踪实际上是一个回报甚微的计划。¹⁰³ 拉尔斯·卡斯特勒 (Lars Kestner) 把趋势追踪系统的绩效与期货投资组合和股票投资组合的绩效进行了比较。期货投资组合考虑了 8 个不同行业的 29 种商品；¹⁰⁴ 股票投资组合考虑的是 9 个行业领域中的 31 只大盘股¹⁰⁵ 以及 3 种股票指数。经过风险调整的绩效 (夏普比率) 计算了 5 种趋势追踪系统，¹⁰⁶ 时间跨度从 1990 年 1 月 1 日至 2001 年 12 月 31 日。夏普比率在期货市场中获得的 5 种趋势追踪系统的平均值与在股票市场中获得的平均值的比率是 0.604 VS. 0.046。这一结果支持了期货市场中的趋势追踪者所赚取的风险溢价是不适用于股票市场中的趋势追踪者的 (见图 7-15)。

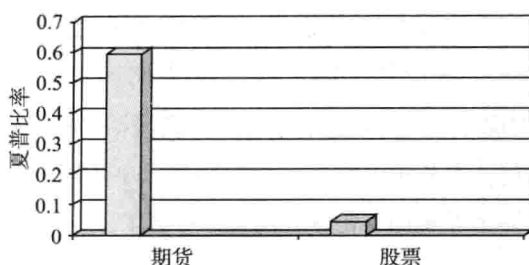


图 7-15 在期货市场中对趋势追踪进行的经过调整的收益率 VS.
在股票市场中对趋势追踪进行的经过调整的收益率

流动性溢价与股票市场中进行反趋势交易获得的利润

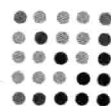
股票市场为接受风险的股票交易者提供了一种不同的补偿形式。他们通过向那些有强烈意愿卖出股票的人提供流动性的方法来赚取溢价。换言之，股票市场中的系统性价格变动可以通过买入最近表现弱势的股票的反趋势策略来实现。

那些急于清算手中股票的股东需要买方。这说明了市场应当向那些愿意满足流动性需求的交易者提供补偿。由迈克尔·库珀提出的实证显示，那些买入超卖股票的买方可以赚取高于短期平均值的收益率。¹⁰⁷ 他的研究表明，在成交量持续下降的情况下，那些表现出负价格惯性的股票是可以赚取超额利润的。换句话说，这表明了那些股票处于亏损状态的卖方正在寻找买方。超额利润可等同于流动性溢价。库珀的研究显示，由于成交量的持续下降，在过去的两周内出现大幅下跌的股票会在接下来的一周内出现上涨的系统性趋势。这种形态中出现的持续下降的交易量之所以看似合情合理，是因为它可以被解释为没有足够的买方来接盘卖方的迫切需求。为了防范数据挖掘的可能性，库珀使用了一种随机游走的样本外模拟，构成了一个从 1978 ~ 1993 年多头头寸与空头头寸的投资组合。这个由 300 个大盘股组成的多头 / 空头头寸所赚取的年化收益率为 44.95%，而采用买入并持有基准获得的年化收益率为 17.91%。我们再次看到了在流动性溢价的情况下，可以用风险溢价解释系统性价格变动，这一点是可以被技术分析所利用的。

在有效股票和期货市场的情况下，由技术分析方法发出的信号可以被视为满足市场需求的好机会。实际上，交易者需要打出这样一则广告，“套期保值者需要雇用风险的接受者，或者急需流动性的提供者——我们将支付工资”。利用技术分析的交易者从这些信号中赚取的利润并不是免费的午餐，他们只不过是读懂了市场的招聘广告罢了。

结语

在本章中，我对金融市场中非随机价格变动的存在做出了解释。如果没有对这些事情的说明，就不可能对技术分析做出合理的解释。在对它们做出解释之后，技术分析就有机会利用非随机行为的某些部分。最后，这全部取决于每一种技术分析方法对自己的证明。



第二部分 Part 2

案例研究：标准普尔 500 指数 的信号法则



第8章

Chapter 8

对应用于标准普尔 500 指数的法则进行数据挖掘的案例研究

本章讲述的是对数据挖掘法则的案例研究，其研究结果我们将在第 9 章中讨论。我们评估了对标准普尔 500 指数进行回溯测试的 6 402 个独立技术分析法则的统计显著性，测试的时间跨度从 1980 年 11 月 1 日到 2005 年 7 月 1 日。

数据挖掘偏差和法则评估

这项研究的主要目的是说明把数据挖掘偏差的影响考虑在内的统计学方法的应用。我们现在回顾一下，数据挖掘实际上是把众多法则的盈利性进行比较之后，挑选出一个或者多个最优法则的过程。正如我们在第 6 章指出的那样，这种选择的过程导致了对被选择法则的绩效出现了过高的偏差。换言之，在回溯测试中观察的最佳法则的绩效夸大了其在未来的预期绩效。这种偏差使得对统计显著性的评估更加复杂，并且导致了数据挖掘器选择了一个没有预测力的法则（以前的绩效完全属于运气）。这就是我们所说的客观技术分析师发现的黄铜矿。

这个问题可以使用特殊的统计推断的方法将其最小化。案例研究对两种方法的应用进行了说明，即怀特真实检验的加强版和重要的蒙特卡罗置换法。它们都利用最近的改进¹降低了忽视（类型 II 错误）优秀法则的概率。也就是说，这种改进增加了测试力。

这项研究的第二个目的是使产生统计显著性利润的一个或者多个法则能够被发现。然而，当进行案例研究的法则被提出时，纵然发现了所有统

计显著性，但对我们来说仍然是未知的。

避免数据探测法偏差

除了数据挖掘偏差以外，对法则进行的研究还遇到了更加严重的问题，即数据探测法偏差。数据探测法偏差是指运用其他研究者记录的对以前的法则的研究结果。由于这些研究通常不会公开数据挖掘的全部数量，所以就无法把它产生的影响考虑进去，从而也就不能对结果的统计显著性做出适当的评估。正如我们在第6章指出的那样，这些都取决于正在使用的是哪种方法——怀特真实检验还是蒙特卡罗置换法，每一种被测试的法则的信息必须都能创建适当的抽样分布。

假设本书的案例研究包括了一种由马丁·茨威格博士的“double 9:1 upside/downside volume”法则。当纽约证券交易所（NYSE）每日上涨和下跌的量比在3个月的时间窗口内超过临界值9的时候，就会发出多头信号。注意，该法则具有3个自由变量：比率的临界值（数值为9）、比率超过临界值的实例数量（2）以及临界值违规的最大时间间隔（3个月）。根据技术分析课堂上学生亲自对该法则进行的测试，信号在接下来的3月、6月、9月以及12月这几个月份中表现出了预测力。换句话说，该法则具有统计显著性，条件是在茨威格不使用数据挖掘来发现他所介绍的参数值9、2和3的假设下才具有显著性。如果当他用不同的参数组合来测试其他版本的法则时，茨威格并没有记录下这种情况，或者说，如果他做了记录，那么又有多少参数组合试图发现对他的法则进行定义的特殊集合呢？如果我们的案例研究包括茨威格法则在内，那么该法则无疑就是最佳的那个了，而我们是可能考虑导致其被发现的确切的数据挖掘数量的。

为了避免数据探测法偏差，我们的案例研究并没有明确包括其他研究者所讨论的法则。尽管有可能，而且是非常有可能出现这种情况，一些用于研究的法则与某些在之前的研究中进行测试的法则是类似的，但是这种类似只是偶然，而不是有意的。这种警惕减轻了数据探测法偏差，但是不能将其完全消除，这是因为，在接受测试的6402个法则当中，我肯定会

受以前提到的法则研究的影响。

分析数据序列

尽管所有法则对它们在标准普尔 500 指数的盈利性做出了测试，但是大多数法则除了利用标准普尔 500 指数以外，还利用数据序列发出买卖信号。其他的数据序列包括：其他市场指数（交通运输行业的股票）、市场宽度（上涨和下跌量）、量价相结合的指标（成交量净额，OBV）、债务工具价格（BAA 级债券）以及利差（利差久期在 10 年期的美国中长期债券与 90 天的短期无息国库券之间）。这种方法以“墨菲论市场互动分析”为基础。²我们接下来将详细讨论所有的序列和法则。

技术分析的主题

接受测试的 6 402 个法则来源于广泛的技术分析主题：①趋势；②极端情况和趋势转换；③背离。趋势法则基于由法则分析的当前数据序列产生多头和空头市场头寸。极端情况和趋势转换法则在所分析的序列达到极端的高值或低值，或者出现在两个极端值之间进行转换的情况时才会产生多头头寸和空头头寸。而背离法则是在标准普尔 500 指数出现同一方向，而与其相伴随的数据序列的趋势却向另一方向发展时，背离法则就会发生信号。有关每一种法则的具体数据转换与信号逻辑问题，我们将在本章后面的内容中描述。

这些法则使用诸如移动平均、通道突破以及随机等一般的分析方法，我把它们称为通道标准化。

绩效统计量：平均收益率

在对处于消除趋势中的标准普尔 500 指数进行回溯测试时，绩效统计量对一定时期内（1980 ~ 2005 年）每一个法则的平均收益率进行了评估。正如我们在第 1 章中指出的那样，消除趋势等于用每天的实际价格变动减

去在回溯测试期间标准普尔 500 指数的平均每日价格变动的差。这将导致一个新数据序列的每日平均变动等于零。

我们在第 1 章中曾经讨论过,消除趋势可以消除任何利润或者消除由于多头或空头偏差引起的法则绩效。如果二进制法则中的一个条件(多头或者空头)相对于另外一个条件(多头或者空头)来说是受到限制的,那么二进制法则就会产生头寸偏差。例如,如果多头头寸比空头头寸更难满足条件,那么该法则就会产生空头偏差,从而就会在相当长的一段时间内持有空头头寸。如果应用于市场数据中的法则具有正趋势的话(平均每日价格变动大于 0),那么其平均收益率就会受到相应的惩罚。法则绩效的降低与其预测力是不相关的,并且无法对评估结果做出正确判断。

不评估复杂法则

为方便对案例研究的范围进行管理,我们只对单一法则进行测试。复杂法则是由两个或两个以上的单一法则用数学或者逻辑运算符合并而成,这些复杂法则我们不予考虑。这种合并的方法既有未加权平均数一般简单,又有从复杂的数据模型软件中推导出来的任意非线性函数一般复杂。

限制对单一法则的案例研究主要有两个原因:首先,少数实践者只依靠单一法则进行决策;其次,复杂法则利用单一法则之间的信息协同效应。因此,复杂法则表现出更好的绩效水平并不意外。有研究³显示,复杂法则可以产生优秀的绩效。即使将简单法则合并成复杂法则也会出现偶尔的不可盈利。合并之后的法则直观上看是非常困难的,但是对合并之后的法则来说,现在有了有效的自动监测方法。⁴

以统计学术语定义的案例研究

下面这部分内容阐述的是统计学关键要素的案例研究:存在争议的总体问题、总体参数、样本、样本统计量、原假设和备择假设以及统计显著性水平。

存在争议的总体问题

存在争议的总体问题是指一组通过应用于标准普尔 500 指数信号获得的每日收益率的集合，而这一信号的应用范围是指马上实现的预期当中能够全部实现的预期。⁵

总体参数

总体参数是指马上实现的预期中法则的预期平均年化收益率。

样本

样本包括应用在消除趋势中的标准普尔 500 指数价格的法则所获得的每日收益率，其回溯测试期从 1980 年 11 月 1 日到 2005 年 7 月 1 日。

样本统计量（检验统计量）

样本统计量是指由法则获得的平均年化收益率，应用在消除趋势中的标准普尔 500 指数的价格数据从 1980 年 11 月 1 日到 2005 年 7 月 1 日。

原假设

原假设阐述的是接受测试的 6 402 个不具有预测力的法则。这说明任何在回溯测试中观察的利润都是偶然的（抽样变异性）。

实际上，在案例研究中存在两种版本的原假设。这种情况要归结于用于评估统计显著性的两种方法：怀特真实检验与蒙特卡罗置换法。尽管两种方法都主张所有的法则都缺少预测力，但是它们对各自假设的应用却是不同的。怀特真实检验认为所有法则的预期收益率等于 0（或者小于 0）；而蒙特卡罗置换法则主张所有法则发出的多头和空头头寸都是随机的。换句话说，多头与空头的排列是随着市场当天的远期价格的变化而变化的。

备择假设

备择假设认为法则在回溯测试中的盈利性来源于其真实的预测力。对

此,上述两种方法的主张稍有不同。根据怀特真实检验,备择假设认为至少有一种在测试集中的法则的预期收益率是大于0的。有一点需要注意的是,怀特真实检验并不认为那个具有最高观察绩效的法则就是最佳法则(拥有最高预期收益率的法则)。然而,在公平合理的条件下,它们就是同一个法则。⁶蒙特卡罗置换法对此的观点是,法则在回溯测试中的盈利性是随着市场当天的远期价格变化而产生的多头和空头信息配对的结果,换言之,该法则是具有预测力的。

统计显著性水平

5%的显著性水平被认作排斥原假设的临界值。这也就意味着当原假设是真实的时候,原假设遭到排斥的概率是0.05。

现实意义

该结果的现实意义与统计显著性之间是存在差异的。统计显著性是指,不具有预测力的法则(原假设是正确的)所获得的收益率等于或者高于在回溯测试中依靠运气获得的收益率的法则。相反,现实意义还与观察的法则收益率的经济价值有关。当样本容量很大时,正如案例研究中的那样(超过6 000天),即使法则的正收益率很小,原假设也会被排斥。换言之,即使法则的实际价值忽略不计,那么该法则还是具有统计显著性的。现实意义取决于交易者的目标:某一交易者也许对某法则获得5%的收益率而感到满足,而其他交易者可能还对低于20%的收益率不满意呢。

法则:将数据序列转换成市场头寸

法则就是一个输入/输出的过程。也就是说,法则把包括一个或者多个时间序列的输入端,转换成为一个包括+1和-1的新输出时间序列,它可以在交易市场中(标准普尔500指数)指出多头和空头头寸。这种转换是通过应用于输入时间序列的一个或者多个数学、⁷逻辑⁸以及时间序

列⁹的运算符来定义的。也就是说, 法则是由一组运算符来定义的(见图 8-1)。

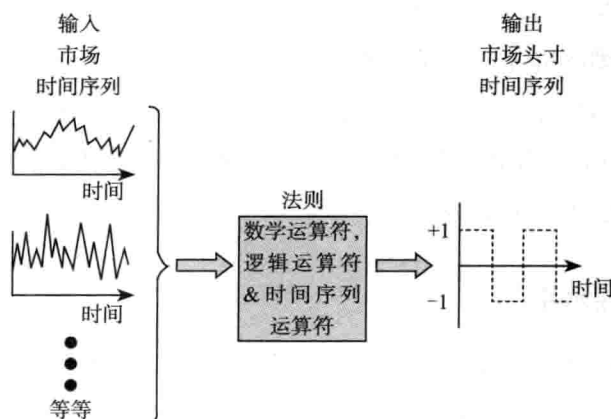


图 8-1 将技术分析法则从输入形态转换为输出形态

有些法则用原始时间序列作为输入端。也就是说, 作为法则输入端的一个或者多个时间序列, 在市场头寸逻辑被应用之前是不能以任何方式进行转换的, 如标准普尔 500 指数的收盘价(见图 8-2)。

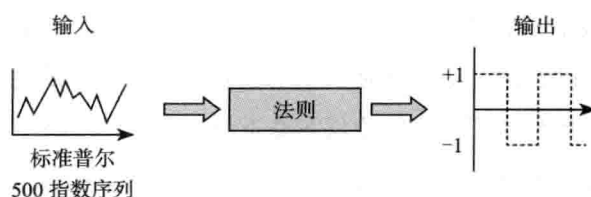


图 8-2 作为法则输入的原始市场时间序列

其他法则的输入序列利用的是来源于一个或者多个的原始市场序列, 这一原始市场序列是通过把各种转换应用到市场数据中获得的。我们把这种经过预处理的输入端称为构成指标的数据系列。这种指标的例子就是负交易量指数(NVI)。它是从对标准普尔 500 指数的收盘价和整个纽约证券交易所的每日成交量这两个原始数据序列的转换中推导而来的(见图 8-3)。对用于构成负成交量指数的转换以及其他用于案例研究的指标, 我们将在接下来的内容中描述。

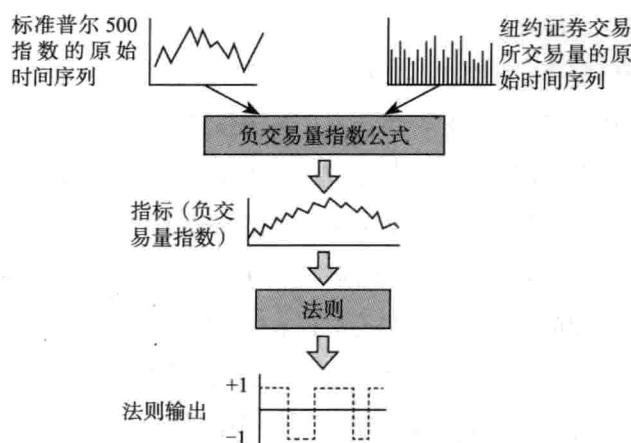


图 8-3 将经过转换的市场时间序列作为法则的输入端

时间序列指标

这部分描述的是，在案例研究中使用的时序指标将原始数据序列转换为指标的内容。时间序列指标是把一组时间序列转换为一组新的时间序列的数学函数。案例研究中使用了几种常见的时序指标：通道突破、移动平均数以及通道规范化。通道突破指标用于确认趋势法则的时间序列趋势。移动平均数指标同样可以用于确认趋势，只是用于平滑趋势。通道规范化指标（随机指标），则应用于消除一组时间序列（消除趋势）的趋势。

通道突破指标 (CBO)

所有趋势确认方法的目的都是审视过去存在于时间序列中的潜在噪声，以及那些看上去毫无意义的不稳定移动，并且发现时间序列当前的趋势。¹⁰ 这样一种方法就是 n 个阶段的通道突破指标。¹¹ 其中， n 代表过去的时期（天、星期等），它用来确认时间序列的上下通道边界。较低的边界由在过去 n 个时期的最小成交量确定，这里面不包括当前的时间段。而较高的边界则是由时间序列在过去 n 个时期内的最大成交量确定，也不包括当前的时间段。尽管这过于简单，但是通道突破指标已经证明了它与复杂趋

势追踪方法¹²同样有效。

按照传统解释, 如果经过分析的时间序列超过了其在过去 n 个时期内确定的最大值, 通道突破指标就会发出多头信号; 相反, 如果时间序列低于其在过去 n 个时期内确定的最小值, 该指标就会发出空头信号, 请参见图 8-4。实际上, 当每个新的数据点变成有效的时候, 这条通道就会被重新画过, 但是图 8-4 中只有几条通道显示有作用。虽然对时间轴来说这条通道有一个恒定的追溯范围, 但是对代表时间序列值的那条轴线来说, 其垂直的宽度却在不断地调整着过去的 n 个时期内时间序列范围。也就是说, 通道突破的动态范围降低了不断增加的时间序列的波动, 而不是趋势中发生的实际变动所导致出现的错误信号的可能性。

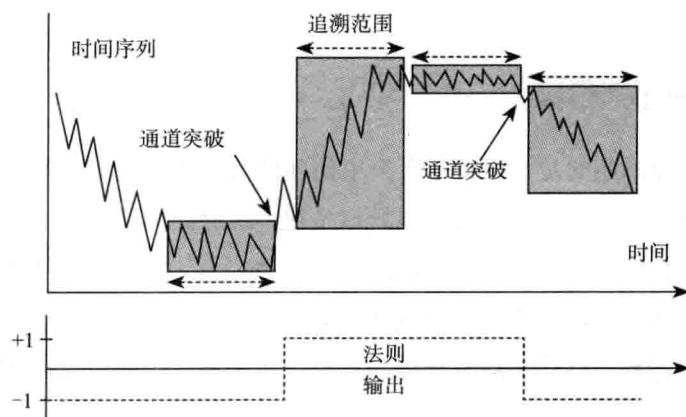


图 8-4 通道突破指标

n 个时期的突破指标具有单一的自由参数, 而其追溯范围、过去时间段的数量确定了通道的边界。总之, 追溯范围越大, 通道越宽, 因此, 其信号的敏感性就比较弱。所以, 较大的数值 n 就会被用来确认大的趋势。我们的研究对 11 种不同的追溯范围进行了测试, 其范围按大约 1.5 倍的长度进行分隔, 具体范围是: 3 天、5 天、8 天、12 天、18 天、27 天、41 天、61 天、91 天、137 天以及 205 天。

移动平均指标 (MA)

移动平均数是技术分析领域中最广泛使用的时间序列指标之一。它通

过在穿过低频率（长期）的部分时过滤掉高频率（短期）波动来说明趋势。因此，移动平均指标被认为具有低通滤波器的作用（见图 8-5）。图 8-5 中，移动平均指标是居中的¹³（延迟 $0.5 \times \text{周期} - 1$ ）。注意，移动平均指标的输出端本质上于消除趋势来说，具有高频率振荡的原始时间序列的趋势。

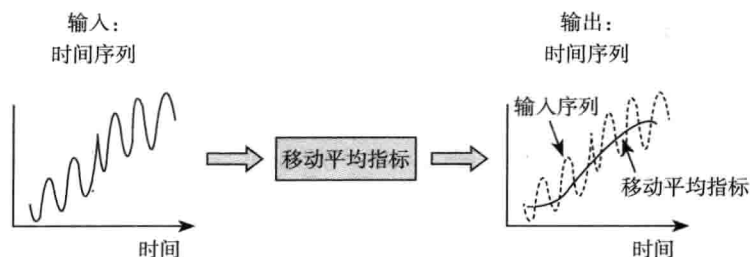


图 8-5 移动平均指标：低通滤波器

平滑作用是通过追溯范围的移动数据窗口内的时间序列的平均值来实现的。平滑作用利用等于或者小于追溯范围期限的平均期限消除或者降低了的波动幅度。因此，对于持续时间小于 10 天的的高峰到高峰（peak-to-peak）或者谷底到谷底（trough-to-trough）的形态来说，10 日移动平均数会降低其波动幅度，并且会完全消除持续 10 天的形态的波动幅度。

然而，要想受益于这种平滑作用就必须付出滞后的代价，换句话说，在平滑的状态中显示出来的原始价格波动确实经历了某些滞后。因此，在原始价格中出现的谷底直到一些时日以后才会在移动平均线中显示出谷底。滞后的时间取决于追溯范围的长度。追溯范围越长，平滑效果越明显，其滞后时间也就越长。

简单移动平均线（simple moving average）的权重作用与在数据窗口中所有因素的权重作用是相等的，由简单移动平均线引发的滞后等于用追溯范围减去 1，然后再除以 2 得出。所以，持续时间为 11 天的简单移动平均线的输出值的滞后期就是： $(11-1)/2$ 或者是 5 天。这意味着，输入时间序列中的趋势反转（其在移动平均线中有足够的时间来证明），将在 5 天后显现。

移动平均线的种类很多，从与在数据窗口中所有因素的权重作用相等的简单移动平均线，再到使用复杂数据加权函数的复合平滑法。复合权重

带来的好处是获得经过改进的过滤绩效, 即减少了对给定平滑程度的滞后程度。数字信号处理考虑的是最大化过滤绩效的权函数。与技术分析相关的数字滤波器可以参考埃勒斯¹⁴的两部著作。

对简单移动平均线的计算可以用下面的公式来表示¹⁵

移动平均指标

$$MA_t = \frac{P_t + P_{t-1} + P_{t-2} \cdots + P_{t-n+1}}{n}$$

$$= \frac{\sum_{i=1}^n P_{t-i+1}}{n}$$

式中, P_t 代表在时间点 t 的价格; MA_t 表示在时间点 t 的移动平均数; n 表示用于计算移动平均数的天数。

简单移动平均线的滞后可以通过使用线性加权移动平均保持平滑程度的方法来降低。正如埃勒斯所指的那样,¹⁶ 他将追溯范围为 n 天 $(n-1)/3$ 的线性加权移动平均线的滞后与简单移动平均线的滞后 $(n-1)/2$ 进行了比较, 因此, 追溯范围为 10 天的线性加权移动平均线的滞后期是 3 天, 而简单移动平均线的滞后期是 4.5 天。线性加权模式对大多数近期的数据点使用的是 n 的加权, 其中, n 表示在移动平均线追溯范围内的天数, 然后再用前一天的权重减 1, 所有权重相加得到的和将作为除数, 请参见下面的等式:

线性加权移动平均

权重

$$WMA_t = \frac{(n) \times P_t + (n-1) \times P_{t-1} + \cdots + (n-(t-n+1)) \times P_{t-n+1}}{n + n-1 + \cdots + n-(t-n-1)}$$

式中, P_t 代表在时间点 t 的价格; WMA_t 代表在时间点 t 的加权移动平均; n 代表用于计算移动平均数的天数。

当指标向上或者向下和重要的临界值交叉的时候, 对应用于发出信号的法则的平滑指标, 案例使用的是 4 天期的线性加权移动平均。平滑形态

减轻了由摇摆指标反复与临界值交叉的不平滑状态所产生的过多信号的问题。有关4天期的线性加权移动平均的计算请参考下面的公式:

4天期的线性加权移动平均

$$WMA_t = \frac{\begin{matrix} \text{权重} \\ \swarrow \quad \downarrow \quad \searrow \quad \nearrow \\ (4) \times P_{t0} + (3) \times P_{t-1} + (2) \times P_{t-2} + (1) \times P_{t-3} \end{matrix}}{10}$$

式中, P_t 代表在时间点 t 的价格; WMA_t 代表在时间点 t 的加权移动平均; n 代表用于计算移动平均数的天数。

通道规范化指标 (随机指数): CN

通道规范化指标 (CN) 消除了在时间序列中的趋势, 因此, 这也说明了围绕在趋势周围的短期波动。通道规范化指标通过衡量在移动通道内的当前位置来消除时间序列的趋势。通道由具体的追溯范围内的时间序列的最大值和最小值所决定。这种方法与之前讨论过的通道突破指标类似。

时间序列的通道规范化形态的设定范围为 $0 \sim 100$ 。当追溯范围内的原始时间序列处于最大值的时候, CN 假设这个值是 100; 当处于追溯范围内的时间序列处于最小值的时候, CN 假设这个值是 0; 数值 50 说明时间序列处于最大值和最小值之间的中间位置。

关于 CN 的计算, 请参考下面的公式。

通道规范化指标

$$CN_t = \left(\frac{S_t - S_{\min-n}}{S_{\max-n} - S_{\min-n}} \right) 100$$

式中, CN_t 代表在第 t 天时第 n 天的通道规范化值; S_t 代表在时间点 t 的时间序列值; $S_{\min-n}$ 表示最后 n 天的时间序列最小值; $S_{\max-n}$ 表示最后 n 天的时间序列最大值; n 表示通道追溯范围的天数。

相关计算的实例如图 8-6 所示。

考虑到滤波特性, CN 指标的作用相当于高通滤波器。与移动平均形成对照的是低通滤波器, 当过滤掉长期波动或者趋势时, CN 则表现出短

期(高频率)振荡。这部分内容如图 8-7 所示。需要注意的是,输入序列具有明显向上的趋势,但是输出序列则没有表现出趋势。

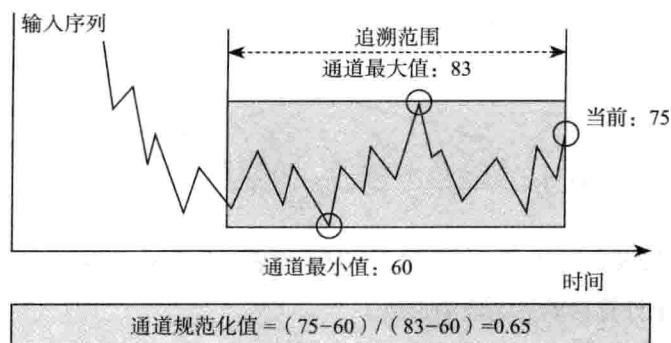


图 8-6 通道规范化

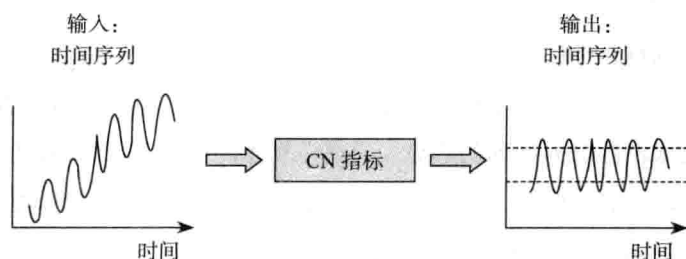


图 8-7 通道规范化形指标: 高通滤波器

CN 运算符已经为技术分析所熟知,并且可以追溯到 1965 年,当时赫比¹⁷利用它把背离法则进行了公式化。他通过对标准普尔 500 指数 CN 值与市场宽度指标的 CN 值进行比较来对背离进行衡量。根据赫比的观点,当一条时间序列处于 50 天区间的上部($CN > 74$),而另一条时间序列则位于该区间的下部($CN < 26$)时,背离就会出现。在 20 世纪 70 年代早期,“随机”这种与 CN 运算符几乎相同的一种运算符由乔治·莱恩(George Lane)¹⁸提出。遗憾的是,这种说法并不准确。实际上,“随机”这个词是一个数学术语,是指时间跨度中随机变量的演变,而对 CN 运算符来说则不存在随机现象,它是一种完全确定的数据转换。本篇所使用的术语是 CN 而不是莱恩提出的术语,这是因为它准确地描述了运算符的含义。大多数技术分析著作仍然在使用术语“随机”。CN 运算符主要适用于两种类型的法则:①极端情况和趋势转换;②背离。我们将在下面的内容

中对其进行讨论。

指标脚本

正如我们之前所讨论的那样，用作法则输入的时间序列要么是原始形态，如来自市场；要么就是指标形态，即通过把一种或者多种数学、逻辑以及时间序列运算符应用到一个或者多个原始时间序列中形成新的时间序列。一个指标可能需要多个运算符的应用。例如，原始时间序列起初可能是一条平滑的移动平均线，然后通道规范化又被应用到这个平滑的移动平均线版本中。当一种指标的产生需要多个步骤时，我们可以很容易地用指标脚本语言（ISL）来表示它。ISL 可以被简单地描述为用于从原始时间序列中产生指标的一组转换。

我接触到的第一种 ISL 的名称是 Gener，它是由我在兰登研究小组的朋友和同事约翰·沃尔伯格（John Wolberg）教授开发出来的。如今，ISL 的使用已经非常普遍。所有的法则回溯测试平台，¹⁹ 如交易大师软件、财富实验室、Neuro-Shell 交易软件以及财务数据计算器等都在使用 ISL。最基础的 ISL 的表述就是，时间序列运算符之后跟上一组用括号括起来的参数。这个参数必须是为运算符所定义的项目。例如，移动平均运算符需要两个参数：要使用运算符的数据序列和在追溯范围内的周期数量。因此，对移动平均指标的表述如下：

移动平均指标 = MA (输入序列, N)

式中，MA 表示移动平均运算符； N 表示移动平均中的天数。

用脚本语言表示的指标的一般形式：

指标 = 运算符 (参数 1, 参数 2...参数 i)

式中， i 表示运算符所需要的参数数量。

图 8-8 显示了到目前为止已确定的运算符的 ISL 排列：通道突破、移动平均以及通道规范化。

ISL 的好处在于其对由嵌套运算符所组成的复杂指标进行定义的能力。在嵌套中，运算符的输入成为第二个运算符的参数，以此类推。例如，如果我们要确定另外一条时间序列（道琼斯运输业指数）10 日移动平均线的

60 日通道规范化指标，那么 ISL 中的指标表示如下：

指标 =CN[移动平均 (道琼斯运输业指数, 10), 60]
通道突破运算符 (CBO)
参数数量 =2
指标 =CBO (时间序列, n)

移动平均运算符 (MA) 或者线性加权移动平均 (LMA)
参数数量 =2
指标 =MA 或者 LMA (时间序列, n)

通道规范化运算符 (CN)
参数数量 =2
指标 =CN (时间序列, n)

图 8-8 关于三种技术分析运算符的指标脚本语言

法则的输入序列：原始时间序列和指标

这部分内容描述的是用于案例研究的 24 个原始时间序列和 39 个最终用作法则输入的时间序列。39 个输入序列中的 10 个是原始的，没有经过处理的时间序列，它们是：标准普尔 500 指数的收盘价、道琼斯运输业指数的收盘价、纳斯达克综合指数的收盘价等；而另外 29 个输入序列是从一个或者多个原始时间序列中推导出来的指标。

原始时间序列

用于案例研究的 24 个原始时间序列的资料来源有两个：超金融体系²⁰和市场时机报告²¹（见表 8-1）。

表 8-1 原始时间序列

序号	原始时间序列	缩写	资料来源
1	S&P 500 Daily Close	SPX	Ultra 8
2	S&P 500 Open	SPO	Ultra 8
3	DJIA High	DJH	Ultra 8
4	DJIA Low	DJL	Ultra 8
5	DJIA Close	DJC	Ultra 8
6	S&P Industrials	SPIA	Market Timing Reports
7	Dow jones Transports	DJTA	Market Timing Reports

(续)

序号	原始时间序列	缩写	资料来源
8	Dow Jones Utilities	DJUA	Market Timing Reports
9	NYSE Financial Indes	NYFA	Market Timing Reports
10	Dow Jones 20 Bonds	DJB	Ultra 8
11	NASDAQ Composite Close	OTC	Ultra 8
12	Value Line Geometric Close	VL	Ultra 8
13	NYSE Advancing Issues	ADV	Ultra 8
14	NYSE Declining Issues	DEC	Ultra 8
15	NYSE Unchanged Issues	UNC	Ultra 8
16	NYSE New 52-Week Highs	NH	Ultra 8
17	NYSE New 52-Week Lows	NL	Ultra 8
18	NYSE Total Volume	TVOL	Ultra 8
19	NYSE Upside Volume	UVOL	Ultra 8
20	NYSE Downside Volume	DVOL	Ultra 8
21	3-Month T-Bill Yields	T3M	Market Timing Reports
22	10-Year T-Bond Yields	T10Y	Market Timing Reports
23	AAA Bond Yields	AAA	Market Timing Reports
24	BAA Bond Yields	BAA	Market Timing Reports

指标

39 个输入中的 29 个是通过用各种各样的数学、逻辑以及时间序列运算符把一条或者多条原始时间序列进行转换之后推导出来的指标。

我们把这些指标分为 4 类：①价量函数；②市场宽度指标；③债务价格工具；④利差指标。

1. 价量函数

有 15 个指标是价量信息相结合的函数。一些研究显示，交易量不仅包括其自身的有用信息，而且还包括与价格一起使用时的信息。²² 技术分析的实践者提出了一些价量函数：能量潮（OBV）、离散量、资金流量、负成交量以及正成交量。我们将在下面的内容中对这些函数逐一介绍。

价量函数可以产生两种类型的指标：①累计额；②移动平均。通过累计额定义的指标是之前每一天的价量函数值的代数和。每天价量函数的数值可以是正数也可以是负数，因此，由 OBV 的累计额定义的指标在某一

个时间点上, 等于之前所有每日 OBV 数值的总和。累计额定义的指标会显示长期趋势, 这与那些观察的资产价格与利率所显示的长期趋势类似。换言之, 这些指标是不稳定的时间序列。

相反, 价量函数的移动平均则是稳定的时间序列。换句话说, 它不显示趋势。这一点可以通过移动平均只考虑在追溯范围之内的观察资料来解释。既然价量函数可以假设为正数和负数, 那么移动平均就会保持在接近 0 的相对有限的范围内 (见图 8-9)。

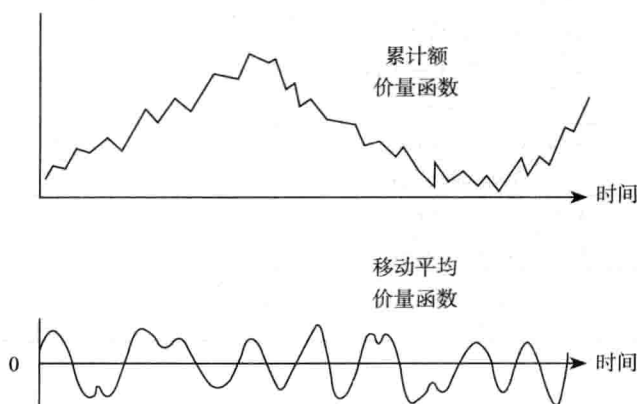


图 8-9 累计额 VS. 价量函数的移动平均

2. 累计能量潮 (OBV)

也许第一个价量函数要属约瑟夫·葛兰碧、伍兹和维格诺里亚发现的累计能量潮指标。²³ 它是一个有正负 (+1 或者 -1) 之分的整个市场交易量的累计额。对于给定日期交易量的代数符号是由最近的某天、某个星期, 或者任意正在使用的时间区间的市场指数的变动决定的。案例研究使用标准普尔 500 指数的每日市场变动来确定适当的交易量符号。如果指定日期的指数出现上涨, 那么当天整个纽约证券交易所的成交量就可以设为正值, 并且把这个值与之前的累计额相加; 如果价格出现下跌时, 整个纽约证券交易所的成交量就可以设为负值, 然后再把这个值加入之前的累计额 (也就是减去)。

关于对累计能量潮 (OBV) 的计算请参见下面的公式²⁴

累计能量潮 (COBV)

$$COBV_t = COBV_{t-1} + (f_t V_t)$$

$$f_t = 1.0 \quad \text{当 } PC_t > PC_{t-1}$$

$$f_t = -1.0 \quad \text{当 } PC_t < PC_{t-1}$$

$$f_t = 0 \quad \text{当 } PC_t = PC_{t-1}$$

式中: OBV_t 代表第 t 天的能量潮; $COBV_t$ 代表第 t 天的累计能量潮; V_t 代表纽约证交所第 t 天的成交量; PC_t 代表第 t 天标准普尔 500 指数的收盘价。

能量潮 (OBV) 的移动平均。除了累计能量潮以外,我们还考虑了另外两种每日能量潮的移动平均,即 10 天和 30 天 (OBV_{10} , OBV_{30})。在 ISL 中,它的表示如下:

$$OBV_{10} = MA(OBV, 10)$$

$$OBV_{30} = MA(OBV, 30)$$

式中, OBV 表示给定日期的能量潮。

3. 累计离散量

马克·柴肯 (Marc Chaiken)²⁵ 提出了另外一种价量函数——累计离散量 (Cumulative Accumulation-Distribution Volume, CADV)。然而,我们还发现了许多类似的函数:大卫·波斯蒂安 (David Bostian) 的每日强度 (intra-day intensity)、拉里·威廉斯 (Larry Williams) 提出的威廉变异离散量 (Variable Accumulation Distribution)。这些变化形式试图克服在 OBV 中出现的感知限制,这种感知限制会把市场指数产生的极小正值(负值)变化认为是全天交易量的正值(负值)。我们认为不应当把较小的价格变动与较大的价格变动相提并论。我们对 OBV 提出修改要求,要么对 OBV 的所有方法进行改进,要么就都不改进。

柴肯提出的每日价格活动特点可以通过对高-低两端范围内收盘价的位置进行量化,以更准确地表现。具体来说,他计算的方法是用收盘价与每日波动范围的中点的差,除以波动范围的差。我们假设,如果市场收盘价在当天的波动范围以上的话,行情上涨或者交易量累积的行为就会在当天的交易活动(累积)中占据优势,并且可以预测价格走高;反过来,当

市场收盘价在当天的波动范围以下的话, 则行情看跌或者交易量分散的行为就会在当天的交易活动中占据优势, 并做出价格走低的预测。只要指数靠近当天交易的低点或者高点, 那么全天的交易量就可以确定为负值(分散)或者正值(累积)。更具体地说, 当指数的收盘价位于当天价格范围内的任何一点时, 就会导致一部分交易量被设定为累积(正值)或者分散(负值)。下面将要介绍的距离因子, 对每日范围内收盘价的位置进行了量化。距离因子乘以每日交易量, 再把这些数字进行累加, 最后得出总和。计算方法请参见下面的公式。²⁶

累计离散量

$$CADV_t = CADV_{t-1} + (V_t \times Rf_t)$$

$$\text{距离因子}(Rf_t) = \frac{(P_{ct} - Pl_t) - (Ph_t - P_{ct})}{(Ph_t - Pl_t)}$$

式中, $ADV_t = V_t \times Rf_t$ 表示第 t 天的离散量; $CADV_t$ 表示第 t 天的累计离散量; V_t 表示第 t 天的纽约证券交易所的成交量; Rf_t 表示第 t 天的距离因子; PC_t 表示第 t 天的道琼斯工业指数的收盘价; Pl_t 表示第 t 天的道琼斯工业指数的最低收盘价; Ph_t 表示第 t 天的道琼斯工业指数的最高收盘价。

案例研究之所以使用道琼斯工业指数, 而不是标准普尔 500 指数用于计算离散函数的数据, 是因为其每日最高价和最低价是无法获得的。我假设在这个距离之内的道琼斯工业指数的收盘位置与标准普尔 500 指数的收盘位置相同。

累计离散量的移动平均。每日离散值的移动平均是由 10 天和 30 天期的离散值组成的, 用 ISL 表示为:

$$ADV10 = MA(ADV, 10)$$

$$ADV30 = MA(ADV, 30)$$

4. 累计资金流量

累计资金流量(CMF)与累计离散量(CADV)除了用交易量和距离因子的乘积再乘以标准普尔 500 指数的价格水平之上, 基本上是相同的。这种指标旨在衡量交易量的货币价值。其计算方法请参见下面的公式²⁷

累计资金流量

$$CMF_t = CMF_{t-1} + (V_t \times Ap_t \times Rf_t)$$

$$\text{距离因子}(Rf_t) = \frac{(P_{Ct} - Pl_t) - (Ph_t - P_{Ct})}{(Ph_t - Pl_t)}$$

式中, MF_t 表示第 t 天的资金流量; CMF_t 表示第 t 天的累计资金流量; V_t 表示第 t 天的纽约证券交易所的成交量; Rf_t 表示第 t 天的距离因子; PC_t 表示第 t 天的道琼斯工业指数的收盘价; Pl_t 表示第 t 天的道琼斯工业指数的最低收盘价; Ph_t 表示第 t 天的道琼斯工业指数的最高收盘价。

资金流量的移动平均。资金流量的移动平均是由 10 天和 30 天移动平均线组成的, 用 ISL 表示为:

$$MF10 = MA(MF, 10)$$

$$MF30 = MA(MF, 30)$$

5. 累计负成交量指数 (CNV)

累计负成交量是指出现最低成交量天数的累积之和。换句话说, 它是指当出现成交量小于前一天的成交量的时候, 这些天股票市场指数的每天百分比变化的代数之和。当成交量等于或者大于前一天的成交量时, CNV 是保持不变的。

我们并不知道 CNV 的发明者是谁, 虽然很多人把它归功于诺曼·福斯贝克 (Norman Fosback),²⁸ 但是也有包括福斯贝克在内的人则认为 CNV 的发明者是保罗·戴萨特 (Paul Dysart)。²⁹ 然而, 可以肯定的是, 福斯贝克开发了适用于 CNV 的客观信号法则, 并且对其预测力进行了评估。³⁰

负成交量的概念基于这样一种猜测, 即以单纯的投资者为主进行的交易出现在交易量上涨时期, 而那些通过消息进行的买卖 (内幕交易) 则出现在成交量下跌时期。因此, 市场假设在负成交量的交易日当天的走势反映了那些拥有内幕消息的人累积 (买入) 或者分散 (卖出) 股票的情况。³¹

该指数的计算方法请参见下面的公式:

累计负成交量指数 (CNV)

$$CNV_t = CNV_{t-1} + f_t$$

$$f_t = \begin{cases} (PC_t/PC_{t-1}) - 1 \times 100 & \text{当 } V_t < V_{t-1} \\ 0 & \text{当 } V_t \geq V_{t-1} \end{cases}$$

式中, $NV_t = f_t$ 表示第 t 天的负成交量指数; CNV_t 表示第 t 天的累计负成交量指数; V_t 表示纽约证券交易所在第 t 天的成交量; PC_t 表示第 t 天的标准普尔 500 指数收盘价。

负成交量指数的移动平均。负成交量指数(指数在成交量较低的那天的百分比变动)的每日价值移动平均是由 10 天和 30 天移动平均线组成的, 用 ISL 表示为:

$$NV10 = MA(NV, 10)$$

$$NV30 = MA(NV, 30)$$

6. 累计正成交量 (CPV)

累计正成交量 (CPV) 与累计负成交量 (CNV) 完全相反。它是指在成交量大于前一天的成交量的情况下, 将这些天市场指数的变化进行代数累计。正成交量指数的提出归功于福斯贝克和戴萨特。³² 其计算方法请参见下列的公式。戴萨特提出了正成交量指数, 并提出了用主观方式对其进行解释的方法, 而福斯贝克则对客观法则以及对它们的正式测试给出了定义

累计正成交量指数 (CPV)

$$CPV_t = CPV_{t-1} + f_t$$

$$f_t = \begin{cases} (PC_t/PC_{t-1}) - 1 \times 100 & \text{当 } V_t > V_{t-1} \\ 0 & \text{当 } V_t \leq V_{t-1} \end{cases}$$

$$f_t = 0 \quad \text{当 } V_t \leq V_{t-1}$$

式中, $PV_t = f_t$ 表示第 t 天的正成交量指数; CPV_t 表示第 t 天的累计正成交量指数; V_t 表示纽约证券交易所在第 t 天的成交量; PC_t 表示第 t 天的标准普尔 500 指数收盘价。

正成交量的移动平均。正成交量每日价值的移动平均是由 10 天和 30 天移动平均线组成的, 用 ISL 表示为:

$$PV10 = MA(PV, 10)$$

$$PV30 = MA(PV, 30)$$

市场宽度指标

市场宽度是指在给定交易日、星期以及其他规定的时间内, 上涨的股票数量与下跌的股票数量的差或者宽度。市场宽度可以通过多种方法来衡量, 我们可以查阅哈罗 (Harlow) 发表的一篇学术研究。³³ 出于研究的目的, 我们通过每日涨跌比指标对其进行定义, 也就是说, 用某一交易日中的上涨值与下跌值之间的差除以已经成交的股票总数。这个数据是以纽约证券交易所中所有已经成交的股票为基础的, 而不是受到限制的普通股的替代版本。对普通股做出限制的时间序列之所以被认为是最佳的, 是因为它排除了封闭式证券、债券基金以及优先股, 而这些证券无法对我们需要的时期以计算机读取的形式获得。

案例研究的市场宽度指标有两种形式: 累计每日价格的累计总额和移动平均。我们之前已经解释过, 通过累计得出的市场宽度指标显示了长期趋势, 而移动平均宽度指标则表现出稳定的平均值与波动范围。

1. 累计涨跌比

累计涨跌比 (CADR) 是每日涨跌比的累计总和。

$$CADR_t = CADR_{t-1} + ADR_t$$
$$ADR_t = \frac{adv_t - dec_t}{adv_t + dec_t + unch_t}$$

式中, $CADR_t$ 表示第 t 天的累计涨跌比; ADR_t 表示第 t 天的涨跌比; Adv_t 表示纽约证券交易所在第 t 天的上涨值; dec_t 表示纽约证券交易所在第 t 天的下跌值; $unch_t$ 表示纽约证券交易所在第 t 天的未变化值。

累计涨跌比的移动平均。每日涨跌比的移动平均是由 10 天和 30 天移动平均线 (ADR10 和 ADR30) 组成的, 用 ISL 表示为:

$$ADV10 = MA (ADV, 10)$$

$$ADV 30 = MA (ADV, 30)$$

2. 累计净成交量比例 (CNVR)

另外一种衡量市场宽度的方法是累计净成交量比例。它是以每日上涨成交量和下跌成交量之间的差为基础的。上涨 (下跌) 成交量是当天交易

的股票的上涨(下跌)数量的总和。我们应当把1938年³⁴对上涨成交量和下跌成交量分别进行计算和分析的创新归功于莱曼·M. 洛瑞(Lyman M. Lowry)。每日净成交量比例的值可以设定为正值或者负值,它是用上涨成交量与下跌成交量之间的差再除以全部的成交量而得出的。本书案例分析使用的统计数据是由纽约证券交易所发布的。累计净成交量比例(CNVR)是每日比例的累计之和

累计净成交量比例(CNVR)

$$CNVR_t = CNVR_{t-1} + NVR_t$$

$$NVR_t = \frac{upvol_t - dnvol_t}{upvol_t + dnvol_t + unchvol_t}$$

式中, $CNVR_t$ 表示第 t 天的累计净成交量比例; NVR_t 表示第 t 天的净成交量比例; $Upvol_t$ 表示纽约证券交易所在第 t 天的上涨值; $dnvol_t$ 表示纽约证券交易所在第 t 天的下跌值; $unchvol_t$ 表示纽约证券交易所在第 t 天的未变化值。

净成交量比例的移动平均。每日净成交量比例的移动平均是由10天和30天移动平均线(NVR10和NVR30)组成的,用ISL表示为

$$NVR10 = MA(NVR, 10)$$

$$NVR30 = MA(NVR, 30)$$

3. 累计新高-低比例

第三种衡量市场宽度的方法是累计新高-低比例(CHLR),它是以长远的观点为基础的。累计涨跌比(CADR)和累计净成交量比例(CNVR)是以每日价格变化和每日成交量为基础的,而累计新高-低比例(CHLR)则是指股票的当前价格,这是相对于股票在过去的一年中的最高和最低价格而言的。具体地说,CHLR是每日新高-低比例(HLR)的累积和。这一数值可以是正值也可以是负值。高-低比例是通过用创造52周新高点那天的股票数量减去创造52周新低点那天的股票数量的差,再除以当天已经成交的数量,我们把这些数值进行累计,就得到了累计新高-低比例(CHLR)。

对使用 52 周这个追溯区间并没有什么特别含义,不过是我们习惯使用这个被广泛使用的统计数据获得新高点和最低点而已。1978 年以前,高点和低点的数值并不是通过 52 周的追溯范围获得的。从任一年份的 3 月中旬起,该数据只以当前的公历年为基础。在 3 月之前,该数据以当年和上一年的数值为基础。因此,在给定年份,用于确定 1 只股票在截至 3 月中旬的这 2.5 个月当中是否创下了新高或新低的追溯范围,其与从上一年度到今年 3 月中旬的这 14.5 个月当中是否创下了新高或新低的追溯范围是一样的。本书的案例研究使用的数据大部分是 1978 年以后的数据。

尽管不同的追溯范围发生了扭曲,但是用 1978 年以及之前的数据进行测试的法则显示,指标的信息内容对这种扭曲却表现出了稳定性。例如,福斯贝克对其开发的一种名为“高-低逻辑指数”的指标在 1944 ~ 1980 年间的表现进行了测试。他的测试结果论证了,即使大量的数据受到了不同追溯范围³⁵造成的扭曲的影响,该指标在这期间仍然具有预测力。组成 CHLR 序列的公式如下³⁶

累计新高-低比例 (CHLR)

$$CHLR_t = CHLR_{t-1} + HLR_t$$

$$HLR_t = \frac{nuhi_t - nulo_t}{adv_t + dec_t + unch_t}$$

式中, $CNHL_t$ 表示第 t 天累计的新高和新低; HLR_t 表示表示第 t 天的新高和新低; $nuhi_t$ 表示第 t 天纽约证券交易所的 52 周高点; $nulo_t$ 表示第 t 天纽约证券交易所的 52 周低点; adv_t 表示第 t 天纽约证券交易所的上涨值; dec_t 表示第 t 天纽约证券交易所的下跌值; $unch_t$ 表示第 t 天纽约证券交易所的不变数值。

新高-低比例 (HLR1 和 HLR30) 的移动平均。每日高-低比例的移动平均是由 10 天和 30 天移动平均线 (HLR10 和 HLR 30) 组成的,用 ISL 表示为

$$HLR\ 10 = MA(HLR, 10)$$

$$HLR\ 30 = MA(HLR, 30)$$

利率中的价格债务工具

利率和股票价格水平的移动通常是相反的。然而，把倒数利率（ $1/\text{利率}$ ）转化成为价格时间序列，在一般情况下，它与股票价格成正相关关系。我们用反转级数乘以比例系数 100，因此，6.05% 的比率就等于价格 15.38 ($1/6.05 \times 100$)。

这种转换将用于案例研究，并且应用于 4 种利率：3 个月期短期无息国库券、10 年期国债、穆迪公司的 AAA 评级公司债券和 BAA 级公司债券。

利差

利差是指两个可以比较的利率之间的差。本书案例使用两种类型的利差：久期利差和质量利差。久期利差也被称为“收益率曲线斜率”，它是指具有相同信贷质量但久期（到期日）不同的两种债务工具的收益率之间的差。应用于案例研究中的久期利差是通过 10 年期中期国库券的收益率减去 3 个月期短期国库券的收益率（10 年期的收益率减去 3 个月期的收益率）来确定的。我们之所以不用 3 年期的收益率减去 10 个月期的收益率，是因为存在于利差之中的上升趋势可能暗示着股票市场（标准普尔 500 指数）会出现牛市。³⁷

质量利差衡量的是久期相同但信贷质量（违约风险）不同的两种债务工具的收益率之间的差。用于案例研究的质量利差是穆迪公司³⁸的两种长期公司债券：最高评级的 AAA³⁹ 级公司债和较低评级的 BAA 级公司债。⁴⁰ 质量利差是通过用 AAA 级债券的收益率减去 BAA 级债券的收益率决定的。当劣质债券（更高的违约风险，如 BAA 级债券）的评级比 AAA 级这类优质债券的评级下降得更快时，我们就可以把这种情况理解为投资者愿意在劣质债券上承担更大的风险，从而赚取更高收益率的标志。换言之，质量利差中的上升趋势是投资者愿意承担更大的风险去赚取更高利润的信号。基于以上几点原因，质量利差的趋势同样适用于股票市场。

应用于案例研究的 40 种输入序列列表

应用于案例研究的 40 种输入序列列表如表 8-2 所示。

表 8-2 应用于案例研究的输入时间序列

序号	类型	简称	形式
1	S&P 500 Close	SPX	Raw
2	S&P 500 Open	SPO	Raw
3	S&P industrials Close	SPIA	Raw
4	Dow Jones Transportation Index	DJTA	Raw
5	Dow Jones Utility Index	DJUA	Raw
6	NYSE Financial Index	NYFA	Raw
7	NASDAQ Composite	OTC	Raw
8	Value Line Geometric Close	VL	Raw
9	Cumulative On-Balance Volume	COBV	Indicator
10	On-Balance Volume 10-Day MA	OBV10	Indicator
11	On-Balance Volume 30-Day MA	OBV30	Indicator
12	Cumulative ASecum. Distr. Volume	CADV	Indicator
13	Accum. Distr. VBvolume 10-Day MA	ADV10	Indicator
14	Accum. Distr. VBvolume 30-Day MA	ADV30	Indicator
15	Cumulative Money Flow	CMF	Indicator
16	Money Flow 10-Day MA	MF10	Indicator
17	Money Flow 30-Day MA	MF30	Indicator
18	Cumulative Negative Volume Index	CNV	Indicator
19	Negative Volume Index 10-Day MA	NV10	Indicator
20	Negative Volume Index 30-Day MA	NV30	Indicator
21	Cumulative Positive Volume Index	CPV	Indicator
22	Positive Volume Index 10-Day MA	PV10	Indicator
23	Positive Volume Index 30-Day MA	PV30	Indicator
24	Cum. Advance Decline Ratio	DADR	Indicator
25	Advance Decline Ratio 10-Day MA	ADR10	Indicator
26	Advance Decline Ratio 30-Day MA	ADR30	Indicator
27	Cum. Up Down Volume Ratio	CUDR	Indicator
28	Up Down Volume Ratio 10-Day MA	UDR10	Indicator
29	Up Down Volume Ratio 30-Day MA	UDR30	Indicator
30	Cum. New Highs/New Lows Ratio	CHLR	Indicator
31	New Highs/New Lows Ration 10-Day MA	HLR10	Indicator
32	New Highs/New Lows Ration 30-Day MA	HLR30	Indicator
33	NYSE Volume	TVOL	Raw
34	Dow Jones 20 Bond Index	DJB	Raw
35	Price 3-Month T-Bill	PT3M	Indicator

(续)

序号	类型	简称	形式
36	Price 10-Year Treasury Bond	PT10Y	Indicator
37	Price AAA Corporate Bonds	PAAA	Indicator
38	Price BAA Corporate Bonds	PBAA	Indicator
39	Duration Spread (10 Year -3 Month)	DURSPD	Indicator
40	Quality Spread (BAA-AAA)	QUALSPD	Indicator

法则

案例研究中接受测试的法则分为3种,每一种都代表着不同的技术分析主题:①趋势;②极端情况和趋势转换;③背离。我们将在接下来的内容里对这些法则进行描述。

趋势法则

趋势法则是以趋势为基础的。技术分析的基本原理就是在趋势中运行的价格和收益可以得到及时确认,并以此获得利润。技术分析的实践者通过开发各种各样的客观技术指标来确定当前趋势的走向,以及发出趋势反转的信号。时下流行的指标有移动平均、移动平均线、通道突破以及有“锯齿型过滤器”之称的亚历山大过滤器等。考夫曼⁴¹对这些指标有过具体的描述,我们在此不再做过多解释。

案例分析使用通道突破指标(CBO)这种趋势法则来确定在输入时间序列中的趋势。因此,CBO指标将输入时间序列转换成为由(+1, -1)组成的二进制值的时间序列。当由CBO确定的输入序列的趋势处于上升趋势时,法则的输出值就是+1;相反,当由CBO确定的输入序列的趋势处于下降趋势时,法则的输出值就是-1。

通过CBO确定的在输入序列中的反转趋势往往具有滞后性。所有的趋势指标都会产生滞后性,即与输入序列所经历的趋势反转的时间相对于指标发现它的时间会产生滞后性。滞后性可以通过让指标更加敏感的方式来降低。就以CBO来举例,我们可以通过减少追溯范围的期限来降低其滞后性。然而,这又产生了另外一个问题,即错误信号数量增加。换句话说,

我们必须在趋势追踪指标的滞后性与其准确性之间做出衡量,因此,所有的趋势指标都试图做到信号的准确性与信号的及时性之间的平衡。不管它到底是家庭烟雾探测器还是技术分析趋势法则,能够找到一个理想的结果对于任何信号系统来说都是一个挑战。最后,我们发现能够把指标绩效最大化的是平均收益率和经过风险调整的收益率等指标。由于金融市场行为时间序列随着时间的推移而改变,对于CBO算子来说最理想的追溯范围也会随之变化。最合适的CBO版本并没有在案例研究中使用。

在表8-2中所列举的40个输入序列中,有39个是作为趋势法则的输入序列的。表中第2项标准普尔500指数的开盘价格由于不属于收盘价格而被排除在外。

每一个趋势法则都通过两个系数或者参数来确定,即CBO的追溯范围以及从表8-2中的39个备选中选取的任意一个时间序列。对追溯范围进行测试的数值由11个值组成,它们分别是3天、5天、8天、12天、18天、27天、41天、61天、91天、137天和205天。我们将每一个选取出来的值分别乘以系数1.5,如205天的追溯范围大约是137天的1.5倍。由此我们获得了858个趋势法则,其中的429个(39×11)是以传统的技术分析说明为基础的,而另外的429个则是反转趋势法则。

当CBO算子确定输入时间序列会出现上涨趋势时,传统的趋势法则版本产生的输出值是+1(标准普尔500指数的多头);根据CBO算子确定时间序列将处于下降趋势时,则输出值为-1(标准普尔500指数的空头)。反转趋势法则只能发出相反的信号(当被分析的序列被确定将会处于上涨趋势时,则做空标准普尔500指数)。

我们在第9章中公布了测试结果,下列的命名惯例将被用于传统的趋势法则和反转的趋势法则。命名惯例对大部分法则来说是非常重要的,其排列顺序如下(注:该惯例与用于指标脚本的惯例是不同的):

法则(TT或者TI)-输入序列数量-追溯范围

其中,TT表示传统的趋势法则;TI表示反转趋势法则。

例如,有一种名为TT-15-137的法则,它应用于表8-2中第15项的使用追溯范围为137天(Cumulative Money Flow)的传统趋势法则。而法则

TT-40-41 则应用于表 8-2 中第 40 个输入序列 (BAA-AAA 的质量利差) 的使用追溯范围为 41 天的反转趋势法则。

极端情况和趋势转换

极端情况和趋势转换法则也称“E 法则”, 它是以当我们假设一个极端高或者极端低的数值, 或者在两个极端值之间进行转换时, 时间序列会传送信息这一概念为基础的。如果时间序列具有一个相对稳定的平均值和稳定的波动范围 (固定的), 那么极端高或者极端低的值就可以通过固定的临界值来确定。所有应用于 E 法则的输入序列通过使用 CN 算子进行固定。

40 个输入序列中的 39 个是适用于 E 法则的。标准普尔 500 指数的开盘价并不包括在内。应用于 E 法则的输入序列在使用 CN 算子之前, 是第一个 4 天期的线性加权移动平均 (LMA) 的平滑序列。如果平滑序列可以降低信号的数量, 那么从信号临界值的上下波动中就可以得出不平滑的输入序列版本。尽管正在进行平滑运动的序列不适用于已经处于平滑运动的序列, 如负成交量的 30 日移动平均, 那么根据一致性, 所有的序列都按照线性加权移动平均来处理。

如果用 ISL 形式的话, 应用于 E 法则的输入时间序列可以用下面的公式来表示:

$CN[\text{线性加权移动平均 (输入序列, 4), 天数}]$

式中, CN 表示通道规范化算子; LMA 表示线性加权移动平均算子。

转换顺序如图 8-10 所示。

当平滑的通道规范化序列与临界值交叉的时候, E 法则就会发出信号。我们假设有较高和较低两个临界值, 当序列与临界值交叉时会产生 (向

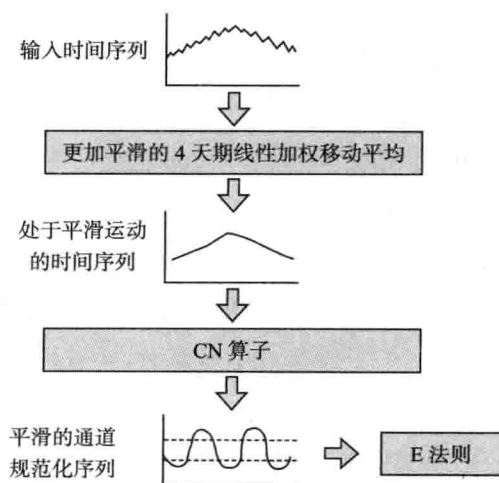


图 8-10 平滑的通道规范化序列

上或向下)两个方向,因此我们就得出了临界值交叉时可能会出现的4种可能。

- (1) 与低临界值交叉产生向下的方向。
- (2) 与低临界值交叉产生向上的方向。
- (3) 与高临界值交叉产生向上的方向。
- (4) 与高临界值交叉产生向下的方向。

请参见图 8-11。

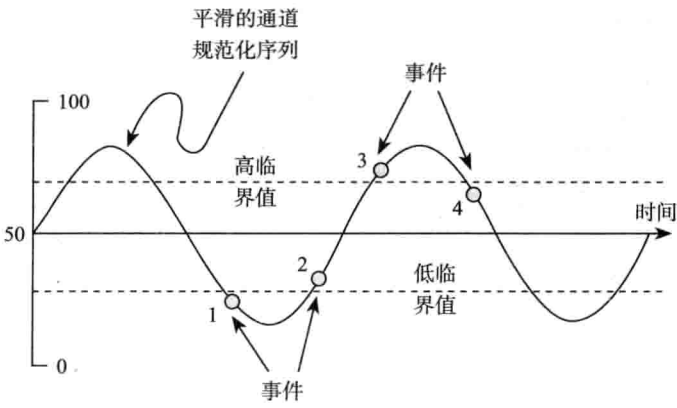


图 8-11 临界值与事件的交叉

每一个 E 法则通过两个临界值与事件进行交叉的方式来确认：多头输入 / 空头输出和空头输入 / 多头输出。这种计划产生出 12 种类型的 E 法则（见表 8-3）。由于表 8-3 包括了所有可能出现的情况，所以就不需要特别指出包括反转类型法则。注意：类型 7 是类型 1 的相反版本，类型 8 是类型 2 的相反版本。

表 8-3 通过临界值与事件的交叉定义的 12 种 E 法则

E 法则类型	多头输入 / 空头输出	空头输入 / 多头输出
1	事件 1	事件 2
2	事件 1	事件 3
3	事件 1	事件 4
4	事件 2	事件 3
5	事件 2	事件 4
6	事件 3	事件 4

(续)

E 法则类型	多头输入 / 空头输出	空头输入 / 多头输出
7	事件 2	事件 1
8	事件 3	事件 1
9	事件 4	事件 1
10	事件 3	事件 2
11	事件 4	事件 2
12	事件 4	事件 3

12 种类型的 E 法则如图 8-12 ~ 图 8-23 所示。

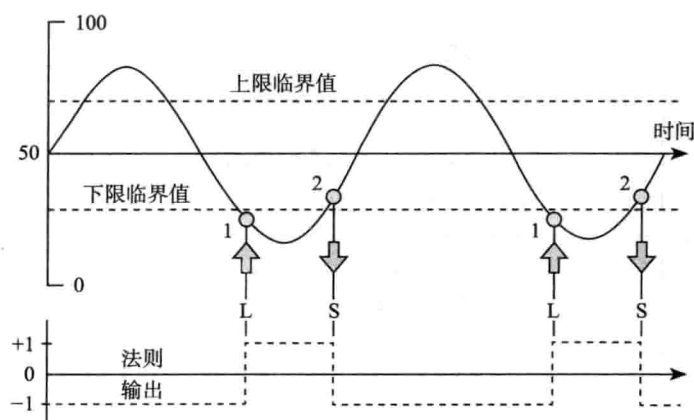


图 8-12 极端情况和趋势转换：类型 1

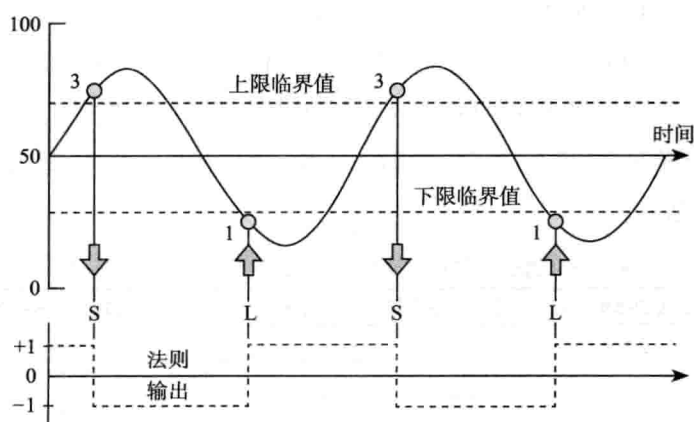


图 8-13 极端情况和趋势转换：类型 2

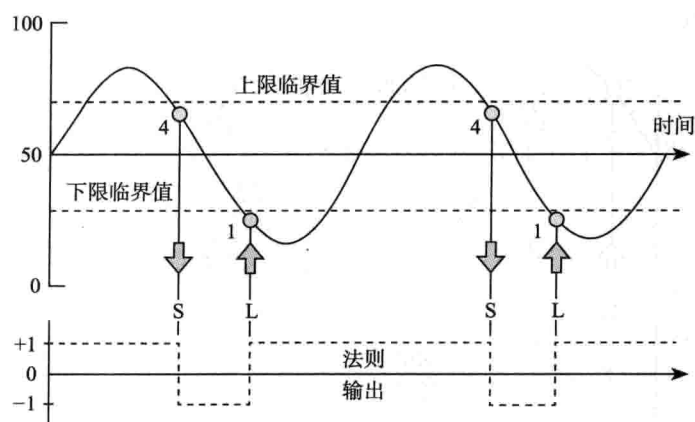


图 8-14 极端情况和趋势转换：类型 3

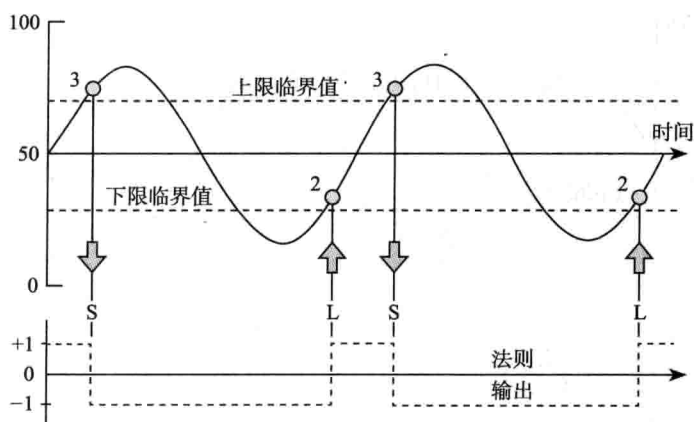


图 8-15 极端情况和趋势转换：类型 4

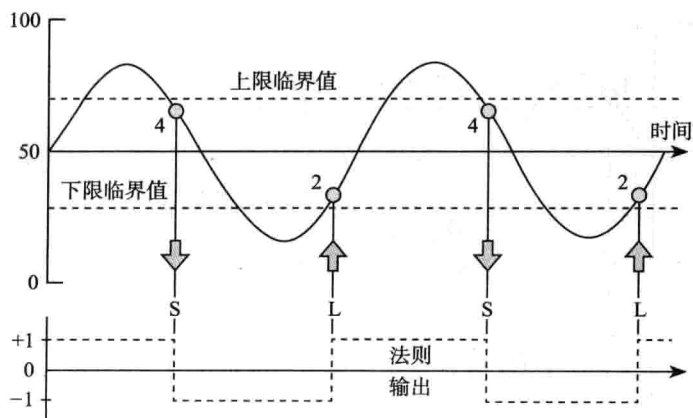


图 8-16 极端情况和趋势转换：类型 5

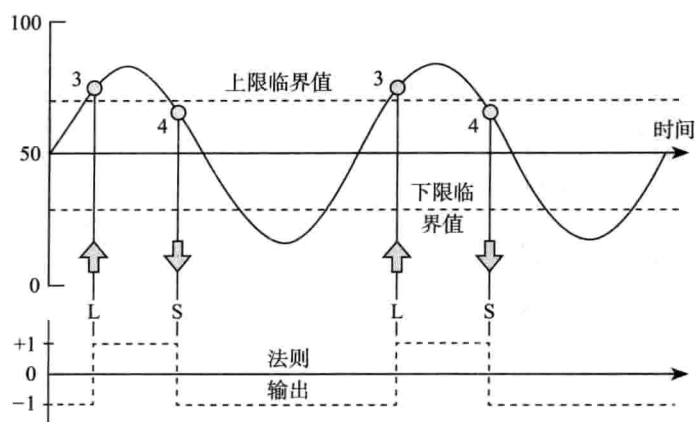


图 8-17 极端情况和趋势转换: 类型 6

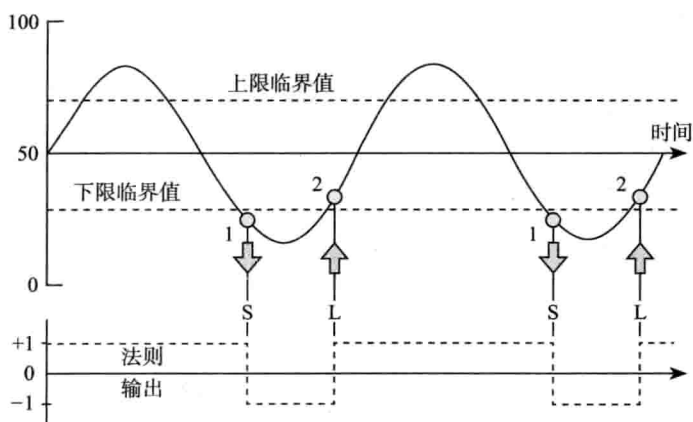


图 8-18 极端情况和趋势转换: 类型 7

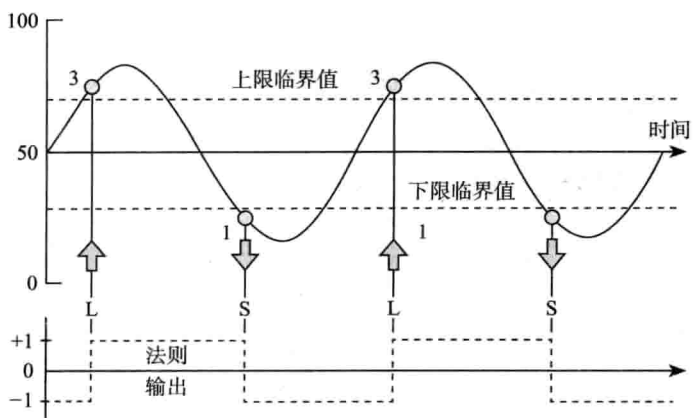


图 8-19 极端情况和趋势转换: 类型 8

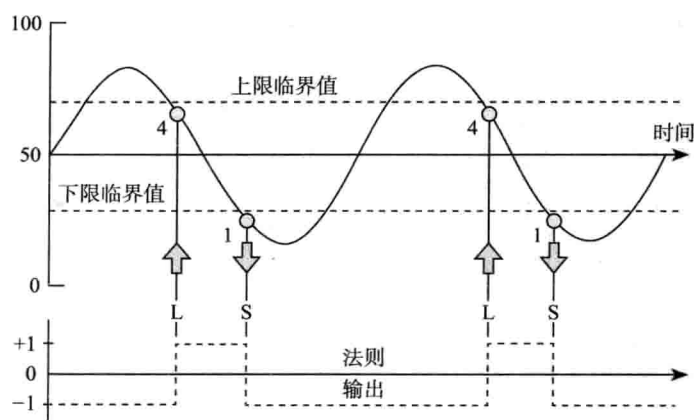


图 8-20 极端情况和趋势转换：类型 9

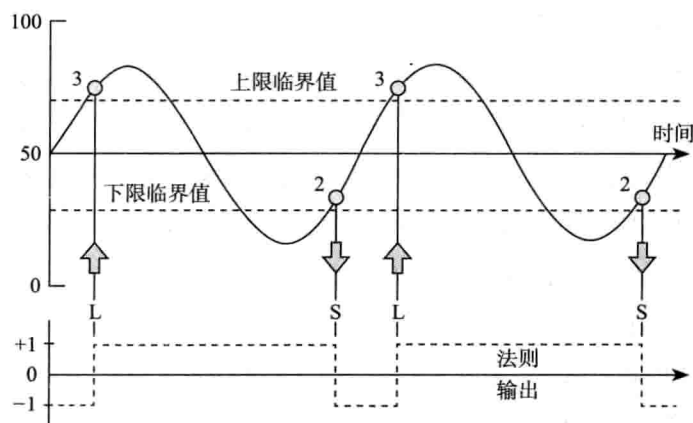


图 8-21 极端情况和趋势转换：类型 10

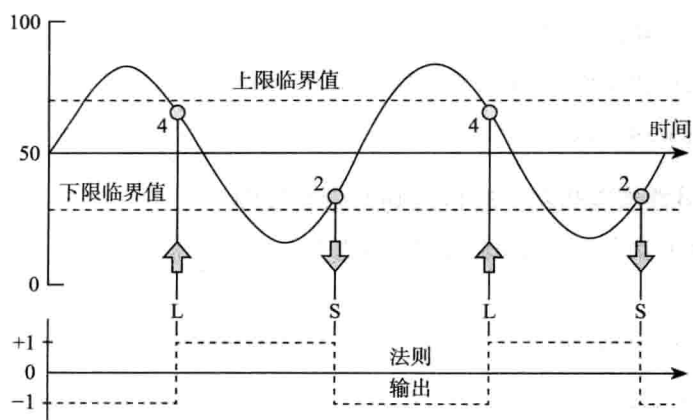


图 8-22 极端情况和趋势转换：类型 11

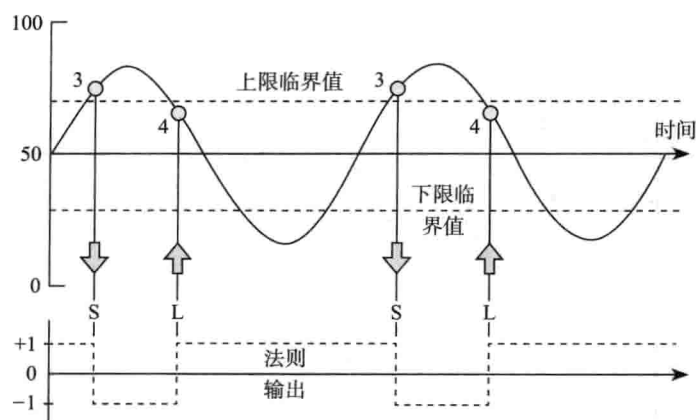


图 8-23 极端情况和趋势转换: 类型 12

1. 参数集和全部 E 法则

E 法则是由 4 个参数确定的: 类型 (类型 1 ~ 类型 12)、输入序列、临界值达到 50 时的位移以及通道规范化的追溯范围。因为上限和下限临界值都是以中值 50 这个点为基础开始出现位移的, 所以我们可以用一个简单的数字来确定上限和下限, 即位移是从数值 50 开始的。例如, 临界值的位移值 10 在上限临界值和下限临界值所处的位置分别是 60 ($50+10$) 和 40 ($50-10$), 因此我们可以计算出两个不同的临界值位移参数值: 10 和 20。而位移值 20 在上限临界值和下限临界值所处的位置分别是 70 和 30。至于通道规范化追溯范围, 我们则使用 3 个不同的数值: 15、30、60。所有的参数值都是我们通过直觉选定的, 并没有经过优化。

我们现在已知 12 种可能出现的 E 法则类型: 39 个备选的输入序列、2 个临界值位移参数、3 个通道规范化的追溯范围, 所以就产生了 2 808 个 E 法则 ($12 \times 39 \times 2 \times 3$)。

2. 极端情况和趋势转换法则的命名规则

在第 9 章中, 下列的命名规则将用来报告 E 法则的结果:

(E) - (类型) - (临界值位移参数) - (通道规范化的追溯范围)

E-4-30-20-60 表示: E 法则、类型 4、输入序列 30 (累计新高 - 低比例)、数值为 20 的临界值位移参数 (上限 70, 下限 30) 以及期限为 60 天的通道规范化的追溯范围。

背离分析

背离分析是技术分析的一个基本概念,它关注的是一组时间序列之间的关系。背离分析的前提条件是:在正常情况下,一组特定的市场时间序列具有共同向上或者向下运行的倾向,并且在走势出现变化时发出相关信息。当两条时间序列中的一条与整个趋势偏离时,背离就产生了。一般来说,背离通过以下方式得到确认:两条序列一直向着同一方向进行延伸,但是,当其中一条序列与其之前的趋势发生了反转,而另外一条序列则依然保持着原来的趋势时,即产生背离。根据背离分析的概念,以上出现的这种情况预示着两条时间序列之前共同保持的趋势出现了减弱的迹象,并且很有可能发生反转,具体情况如图 8-24 所示。赫斯曼认为,当我们对大量的时间序列进行分析后发现其中一些序列开始出现偏离时,此时的背离信号最具信息含量。⁴²

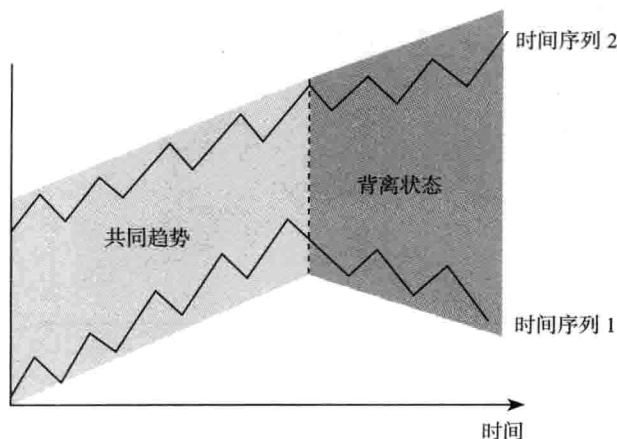


图 8-24 背离分析

道氏理论就是以背离分析为基础的。它认为,当工业类股票指数和交通运输类股票指数的趋势处于同一方向时,两种指数的趋势是健康且容易持续的。然而,如果其中的一条序列出现偏离,我们就可以初步判定该趋势正在减弱并且有可能出现反转。背离分析的另外一种应用是,把一种投资工具的价格看作第一条时间序列,而把其变动率或者趋势看作第二条时间序列。普林格⁴³对价格/趋势背离分析进行了论述。

背离可能会导致出现两种结果中的一种：要么保持趋势不变的时间序列发生反转并加入到另一条发生反转的序列趋势当中，要么正在发生背离的序列终止其错误的趋势并且重新回到与另一条序列的共同趋势当中。直到两条序列出现令人信服的反转并重新向同一方向运动时，最完整的反转信号才会被认为是真实的。因此，背离分析的基本思想就是一致性，也就是说，两条波形的形态是一致的（见图 8-25）。当两条序列形态一致时，它们产生的共同趋势就可以被认为是强势且持续的。然而，当两条序列出现不一致时，或者是不协调时，这种共同趋势的未来走势就会出现問題。因此，作为背离分析备选的序列就是那些通常可以产生一致的时间序列，但是，当出现背离时，这组序列中的一条序列就会有主导另一条序列的信息。换言之，就是两者间具有一种稳定的超前—滞后关系。

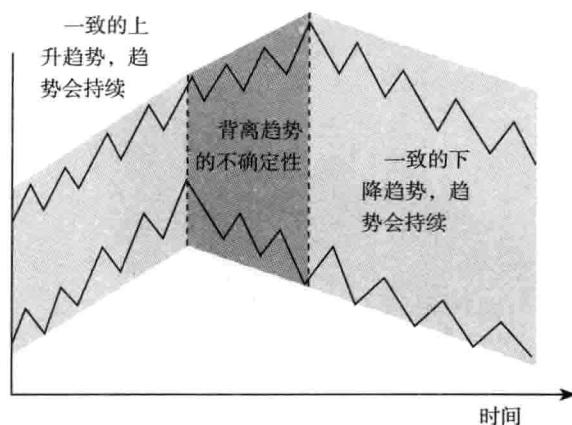


图 8-25 趋势一致和背离

1. 主观的背离分析

主观的背离分析包括对两条时间序列的顶部和谷底进行比较。如果一条序列在连续的高位创出高点，而另外一条序列却在更低的水平上形成顶部时，负的或者是熊市背离（bearish divergence）就会产生。如果第二条序列没有在连续的高位形成顶部时，那么就会出现熊市偏差（见图 8-26）。

当一条序列在已经确立下降的趋势中连续的创出低点，而另外一条序列却开始在高位形成谷底时，就会出现正的或者牛市背离（bullish divergence），我们把这种情况称为牛市偏差（bullish nonconfirmation）（见

图 8-27)。

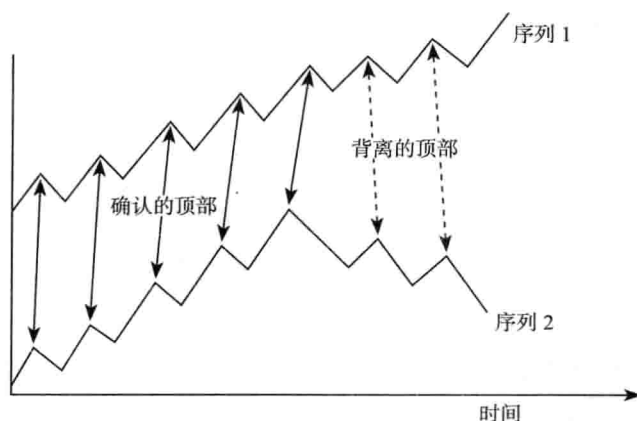


图 8-26 负背离 (对顶部进行比较)

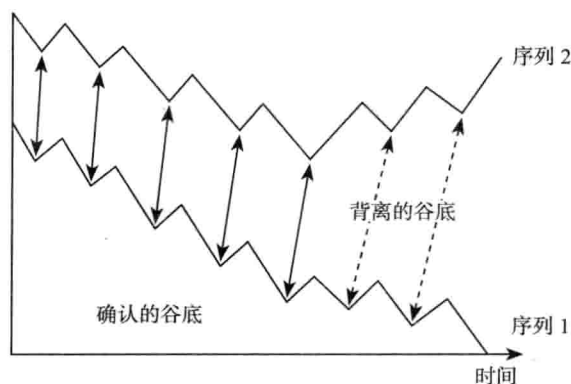


图 8-27 正 (牛市) 背离: 对谷底的比较

主观的背离分析或者说主观的方法所带来的问题就是,这一过程既不可重复,又不可检测。首先,我们要求分析师对顶部和谷底的时间做出主观判断,这是因为两条时间序列不可能正好在同一时间出现,所以也就不好判断应该对哪个顶部和谷底进行比较;其次,我们要考虑久期的因素:在确定哪条序列出现背离之前,我们必须知道背离已经出现多久了;最后,在一条序列当中,需要用多少上升(下跌)的移动来建立新的谷底(顶部)。

2. 对背离进行客观的衡量

上面提到的这些问题可以用客观的背离分析法来解决。考夫曼⁴⁴把进

行比较的顶部和谷底具体化, 并且提出了计算机代码。考夫曼还提出了第二种方法, 即利用线性回归来评估两条时间序列的斜度。背离是通过两条序列之间的斜度差异确认的。这两种方法听起来非常有趣, 但是考夫曼没有提供任何绩效统计数据。

案例研究中所考虑的背离法则受到了赫比⁴⁵提出的客观的背离方法的影响, 赫比在其著作 *Stock Market Profits through Dynamic Synthesis* 中对此进行了论述。赫比利用通道规范化对其分析的两条序列第一次实行了消除趋势的操作。如果一条序列的通道规范化值等于或者大于 75, 而另外一条序列的值小于 25 的话, 背离就会得到确认。因此, 我们通过不同的通道规范化值对背离进行了量化。

我们在下面的公式中提出了对背离指标的确认。请注意, 其中的一条序列通常是指标准普尔 500 指数, 也就是案例研究中的目标市场; 另一条序列是道琼斯交通运输指数。案例研究与标准普尔 500 指数进行比较的所有时间序列 (标准普尔 500 指数的开盘价并不包含在内), 如表 8-2 所示。

背离指标 (初始公式)

$$=CN(\text{比较序列}, n) - CN(\text{标普 500}, n)$$

式中, CN 代表通道规范化算子; n 代表通道规范化的追溯范围。

因为每一条序列的通道规范化值为 $0 \sim 100$, 所以背离指标的潜在范围就是从 $-100 \sim +100$ 。例如, 如果其中一条序列在通道范围的底部 ($CN=0$), 而标准普尔 500 指数则位于该范围的顶部, $CN=100$, 那么标准普尔 500 指数的背离指标值就是 -100 。需要注意的是, 我们这里对背离进行的量化仅仅是众多优秀方法中的一种。

3. 背离指标的局限性

背离指标衡量的是在各自的通道内具有相同位置的两条时间序列的强度。一般来说, 该指标将两条一致序列的强度进行了合理量化。例如, 当背离指标的数值为 0 时, 则表示目前没有出现背离, 我们可以推测两条具有相同的通道规范化值的序列会保持共同的趋势。然而, 也会出现当背离指标的数值为 0 时, 两条时间序列并没有出现相一致的情况。如果两条时间序列呈负相关关系, 也就是说, 这两条序列在通道内的走势是相反的,

而背离指标将会把数值0视为与通道交叉的标准化序列值。在这种情况下,数值0就是预测两条时间序列具有相同趋势的错误指标。这种情况如图8-28所示,我们可以清楚地看到背离指标的限制性。

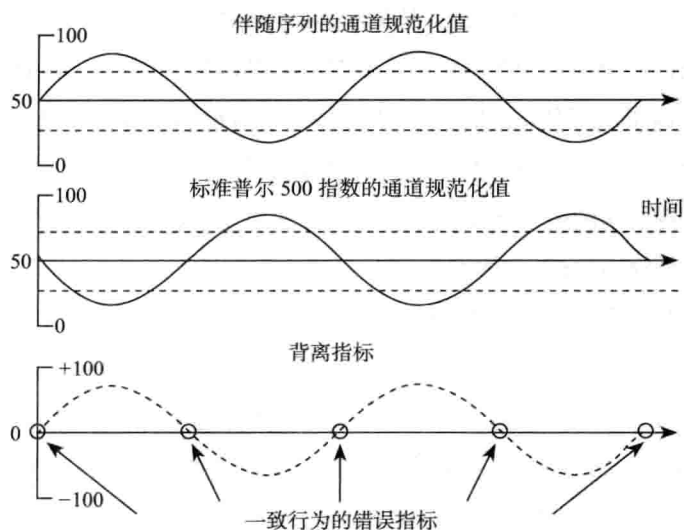


图 8-28 一致行为的错误指标

这个问题可以通过在趋同性⁴⁶的基础上建立复杂的背离指标公式的方法得到解决。1987年,恩格尔(Engle)和格兰杰(Granger)⁴⁷提出了趋同性的概念,这一概念在计量经济学领域中得到了发展。这种更复杂也更准确的衡量背离的方法利用回归分析来确定两条时间序列之间是否存在线性关系。因为这种方法更适用于具有强烈负相关关系(180度反向)序列的情况,因此它解决了我们之前讨论过的在通道规范化的基础上出现的背离指标的问题。趋同性分析的另外一个优势是,它是第一个使用统计检验⁴⁸来确定两条序列是否具有相关趋势的(即两条序列是否存在趋同性)。如果检验的结果显示它们具有趋同性,那么稳定的背离指标⁴⁹就会被分析法所排除。如果该指标使用趋同性的术语来解释,就是“误差校正模型”(error-correction model)。该模型衡量的是一条序列与两条序列之间的线性关系相背离的程度。

在大多数的趋同性应用当中,误差校正模型用于预测趋同性时间序列之间的扩展行为。在这些应用中,背离的作用是在误差校正模型出现极端

的正值或者负值时,恢复其共同趋势的正常形态。换言之,当错误(与两条序列之间的线性关系出现背离)出现极端的数值时,背离就会把它修正到正常数值0的水平。

“配对交易”(pairs trading)是关于这种分析的典型应用。如果我们已知福特公司的股票与通用公司的股票是趋同性的时间序列,且福特公司的时间序列要高于通用公司的时间序列,这时我们就会对福特公司的股票持有空头头寸,而对通用公司的股票持有多头头寸。因此,不论对背离做何种修正(福特公司的股票下跌,通用公司的股票上涨,或者是两种情况的混合),配对交易都会产生利润。

福斯贝克⁵⁰对趋同性分析的使用进行了创新,并且开发出了股票市场预测器。他的创新把从误差校正模型中推导出来的测试结果作为一种指标,来预测第3个变量(即标准普尔500指数,而不是预测两个趋同性变量之间的差额行为)的变化,我们把这个指标称为“福斯贝克指数”(Fosback index)。它利用的是利率和共同基金现金水平是趋同性时间序列的这一事实。福斯贝克的研究之所以引起了特别的关注,是因为他的研究要比恩格尔和格兰杰对趋同性的研究早10多年。当共同基金的现金水平与由它和短期利率的线性关系所预测的水平出现显著背离时,福斯贝克指数就会发出信号。而发出信号的前提就是,当共同基金的经理对股票市场的前景过度悲观的时候,他们的资金储备要比通过利率估计获得的资金多得多。过度的乐观则是另外一种情况。福斯贝克发现通过这种方式对过度悲观和过度乐观进行衡量,与未来股票市场的收益率相关联。实际上,福斯贝克指数把短期利率的波动从共同基金的现金水平中移除,因此,对基金经理情绪的衡量要比通过受到利率影响的原始现金水平进行的衡量更加绝对。

基于简化的考虑,趋同性技术并没有用于案例研究的背离指标之中。我们假设伴随序列与标准普尔500指数无关,那么它就会通过背离法则糟糕的绩效评价显示出来。⁵¹我相信趋同性方法论对指标发展的进一步研究是非常有必要的。

4. 对双重通道规范化的需求

背离指标的初始版本产生的第2个问题更加严肃且必须解决。在案例

研究中的背离法则用标准普尔 500 指数与 38 条时间序列进行配对。假设这些序列与标准普尔 500 指数具有不同程度的联动，那么每组配对的背离指标的波动范围就会出现很大的变动，这对所有的配对组合使用相同的临界值来说是不切实际的。关于这个问题如图 8-29 所示。请注意，适用于与标准普尔 500 指数联动性低的伴随序列的临界值位移，是永远不会对与标准普尔 500 指数联动性高的伴随序列发出信号的。出于以上原因，背离指标的初始公式是不切实际的。

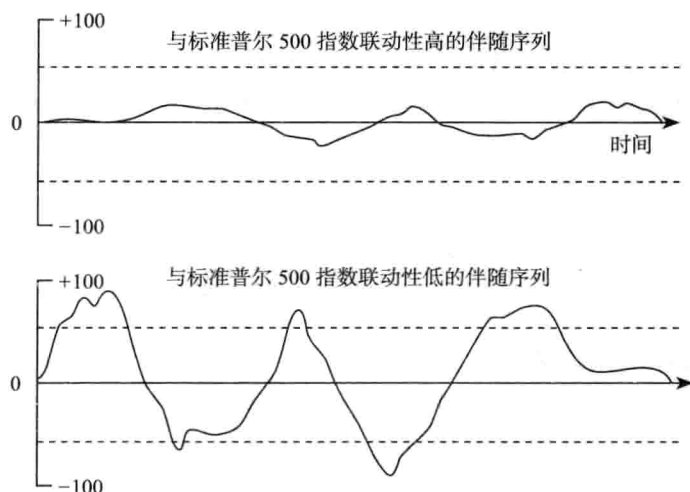


图 8-29 背离指标：波动的不一致性

我们通过对下面背离指标公式的修正来处理这个问题。该方法使用了 2 次通道规范化算子，也就是说，指标就是初始背离指标的通道规范化版本。

背离指标

(双重通道规范化)

$$=CN\{CN(\text{序列 } 1, n) - CN(\text{标普 } 500, n), 10n\}$$

式中， CN 表示通道规范化算子；序列 1 表示伴随序列； n 表示第 1 个通道规范化的追溯范围。

通道规范化的第二个层面考虑的是背离指标初始公式的波动范围。这种方法解决了 38 种序列组合之间不一致的波动范围。结果是，由于经过修正的背离指标会具有相同的波动范围，所以我们需要使用统一的临界值，

而不用去考虑正在使用的时间序列。

通道规范化第二个层面的追溯范围是其第一层面使用的追溯区间的 10 倍。也就是说, 如果通道规范化的第一层面使用的追溯范围是 60 天的话, 那么第二层面的使用范围就应该是 600 天, 我们假设 10 倍的追溯范围足以建立基本背离指标的波动范围。请注意, 经过修正的背离指标的潜在波动范围是从 0 ~ 100, 这与任何通道规范化变量都类似 (见图 8-30)。

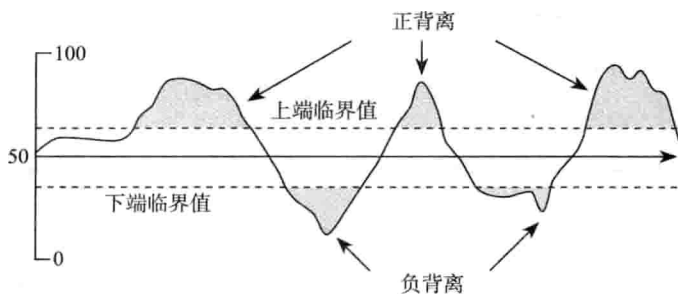


图 8-30 经过修正的背离指标

5. 背离法则的类型

我们把上端临界值和下端临界值应用到经过修正的背离指标中, 让其发出信号。当背离指标超过上端临界值时, 就会产生正的或者牛市背离。当伴随序列的向上移动高于或者向下移动低于标准普尔 500 指数的时候, 这种情况就会发生。我们通过伴随序列在其通道内的相对位置高于校准普尔 500 指数的事实对此进行了证明。相反, 当背离指标低于下端临界值时, 就会产生负的或者熊市背离。当伴随序列的向下移动高于或者向上移动低于标准普尔 500 指数的时候, 这种情况就会发生。结果是, 伴随序列在其通道内的相对位置要比标准普尔 500 指数低 (见图 8-30)。

但问题是, 我们如何从经过修正的背离指标中生成信号法则呢? 有两个法则是显而易见的: 牛市背离法则和熊市背离法则。当背离指标高于上端临界值时, 牛市背离法则就要求对标准普尔 500 指数持有多头头寸; 而当背离指标低于下端临界值时, 牛市背离法则就要求对标准普尔 500 指数持有空头头寸。

两种法则都假设伴随序列拥有主导标准普尔 500 指数的信息。如果伴

随序列强势于标准普尔 500 指数（正背离），那么多头头寸就会得到确认；而如果伴随序列弱势于标准普尔 500 指数（负背离），那么空头头寸就会得到确认。然而，这种假设也许并不十分正确。也有可能会出现其他情况，即负背离具有预测标准普尔 500 指数价格走高的能力，而正背离具有预测标准普尔 500 指数价格走低的能力。换言之，有可能只有标准普尔 500 指数自身才具有主导信息。这种原因就提示我们应该对每一种法则的相反版本做出尝试。最后，我们确定有 12 种可能出现的背离法则类型，并对其逐一进行测试。

用于极端情况和趋势转换法则的这 12 种法则都是相同的集合。由于经过修正的背离指标与用于 E 法则的指标类似，其波动范围都是从 0 ~ 100，且具有 2 个临界值，所以这种说法是合理的。

表 8-4 中所列的 12 种背离法则类型，基本的牛市背离法则（类型 6）、熊市背离法则（类型 7）以及它们的相反版本（类型 12 和类型 1）。要把 12 种类型的法则在极端情况和趋势转换法则这部分内容中都予以描述的话会显得非常冗余，所以在此我们就不逐一说明了。

表 8-4 背离法则与相关的临界值事件

背离法则的类型	多头输入 / 空头输出	空头输入 / 多头输出
1	Down cross Lower Threshold Down Cross	Upper Cross Lower Threshold Upper Cross
2	Lower Threshold Down Cross	Upper Threshold Down Cross
3	Lower Threshold Up Cross	Upper Threshold Up Cross
4	Lower Threshold Up Cross	Upper Threshold Down Cross
5	Lower Threshold Up Cross	Upper Threshold Down Cross
6	Upper Threshold Up Cross	Upper Threshold Down Cross
7	Lower Threshold Up Cross	Lower Threshold Down Cross
8	Upper Threshold Down Cross	Lower Threshold Down Cross

(续)

背离法则的类型	多头输入 / 空头输出	空头输入 / 多头输出
9	Upper Threshold Up Cross	Lower Threshold up Cross
10	Upper Threshold Down Cross	Lower Threshold up Cross
11	Upper Threshold Down Cross	Lower Threshold up Cross
12	Upper Threshold	Upper Threshold

6. 参数组合与背离法则的命名规则

每一个背离法则都由 4 个参数定义：类型、伴随序列、临界值位移以及通道规范化追溯范围。现在有 12 种背离法则（见表 8-4）、38 条数据序列、2 个临界值位移数值——10 和 20，以及 3 个追溯范围——15 天、30 天和 60 天，因此我们就得到了总共 2 736 个背离法则（ $12 \times 38 \times 2 \times 3$ ）。

用于报告背离法则和第 9 章提到的背离法则结果的命名规则如下：（背离法则）- 类型 - 伴随序列 - 临界值位移 - 通道规范化追溯范围。因此，法则的命名如下

D-3-23-10-30

其依次表示：背离法则、类型 3、伴随序列 23（正成交量指数 30 天移动平均）、临界值位移为 10（上限临界值为 60，下限临界值为 40）、通道规范化追溯范围为 30 天。

我们对所有在案例研究中进行测试的法则进行了完整的阐述，其结果请参见第 9 章。

第9章 Chapter 9

案例研究结果与技术分析的未来

研究结果

案例研究的第1个目的是，论证适用于通过数据挖掘发现法则的两种统计学推断方法的应用。正如我们在第6章中解释的那样，传统的显著性检验并不十分合适，这是因为它们没有把数据挖掘的偏差效应考虑进去。我们所用的两种方法分别是怀特真实检验和蒙特卡罗置换法。

案例研究的第2个目的是，当对标准普尔500指数进行研究时发现具有统计显著性收益率的法则的可能性。为了达成这个目的，我们将对在第8章中描述的那6402个法则的集合进行回溯测试和评估。

至于主要目的，案例研究已经成功地论证了使用旨在处理数据挖掘偏差的显著性检验的重要性。关于案例研究的第2个目的，我们没有发现一个具有统计显著性收益率的法则。具体来说，在全部的6402个法则当中，没有一个法则的回溯测试平均收益率在0.05的显著性水平上，可以达到拒绝原假设的程度。换言之，实证并不足以拒绝没有一个法则具有预测力的假设。

拥有最佳绩效的法则E-12-28-10-30，¹在消除市场数据的情况下，产生了10.25%的平均年收益率。在图9-1中，我们把法则的收益率与由怀特真实检验产生的抽样分布进行了比较。在图9-1中，收益率的概率值是0.8164，远远超过了设定值为0.05水平的显著性临界值。图9-1表明了普通的抽样变异性范围内，所有被考察的6402个法则当中的那个最佳法则的绩效同样出现了下降，实际上，它已经跌破了抽样分布的中心值11%。

这一点太令人失望了! ²

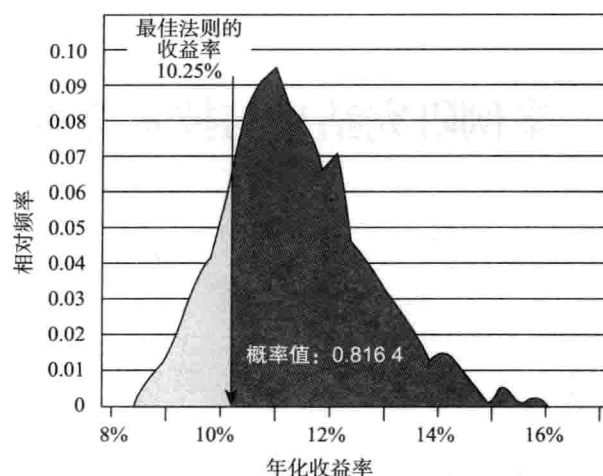


图 9-1 在使用怀特真实检验方法的 6 402 个毫无价值的法则中,
那个表现最佳的法则的抽样分布

上面给出的抽样分布是由 1 999 个反复进行的自举过程产生的。被反复运用的数量越大, 分布的形态就越平稳, 但是有可能会出现相同的结果: 法则的绩效并不能够达到足够的高度来拒绝原假设。

图 9-2 显示了法则 E-12-28-10-30 在由蒙特卡罗置换法产生的抽样分布中的情况。它的概率值是 0.819 4。

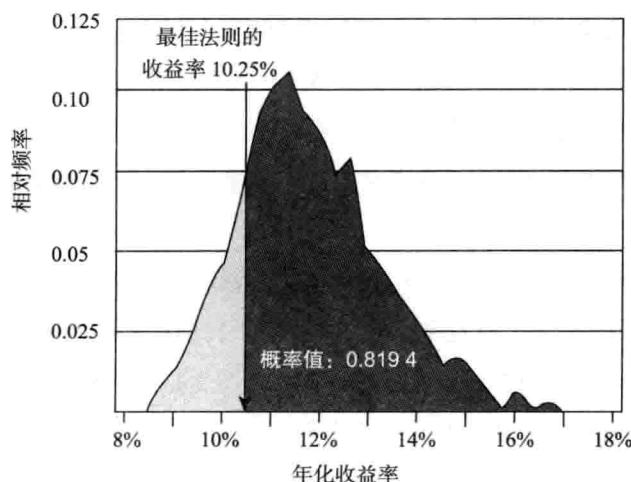


图 9-2 在使用蒙特卡罗置换法的 6 402 个毫无价值的法则中,
那个表现最佳的法则的抽样分布

请注意,这两个抽样分布的平均收益率都集中在 11% 的位置上。这个特殊的数值就是我们的研究结果: 6 402 个法则的集合, 收益率的相互关联性, 用于计算它们平均收益率的观察资料的数量, 以及特定的标准普尔 500 指数的数据集合。换句话说, 在 11% 这个特定的条件下, 将最好的 6 402 个法则的预期收益率与不具有预测力的法则进行比较, 结果不是 0。

我们再来回顾一下第 6 章中提到的内容, 在 N ($N = 6\,402$) 种法则中产生最佳绩效法则的平均收益率的抽样分布并不是以 0 为中心的, 准确地说, 它位于反应统计量抽样分布 (定义为 N 个平均值的最大平均值) 的正值这个点上。

图 9-1 和图 9-2 给出了达到统计显著性所需要的收益率的提示。在 0.05 水平上出现的显著性所需要的收益率为超过 15%, 收益率超过 17% 就是高度显著性 (概率值小于 0.001)。

具有讽刺意味的是, 在经过数据挖掘偏差调整后的法则之所以没有产生统计显著性的收益率, 是为了强调统计学推理方法的使用 (该方法把数据挖掘的偏差效应考虑了进去) 的重要性。如果我使用的是普通的显著性检验, 而没有把数据挖掘偏差考虑进去, 那么最佳法则的平均收益率就会显示出很高的显著性 (概率值为 0.000 5) (见图 9-3)。图 9-3 显示了适用于对单一法则进行回溯测试 (忽略数据挖掘偏差) 的自举抽样分布。与图 9-1 图 9-2 相反的是, 图 9-3 中的抽样分布集中在数值 0 处。图 9-3 中的箭头代表最佳法则 (E-12-28-10-30) 的平均收益率。

如果我们对 6 402 个法则中的 320 个法则使用了传统的显著性检验的话, 那么其显著性就会出现在 0.05 的水平上。这就是我们预测会偶然发生的事情。而使用传统显著性检验方法的天真的数据挖掘者, 就会得出他们发现了很多具有预测力的法则的结论。实际上, 用这种方法进行数据挖掘只会挖到“傻瓜的黄铜矿”。

对全部法则集合进行测试得出的结果可以在 www.evidencebasedta.com 上获取。我们把其中产生平均收益率最高的 100 个法则用表 9-1 表示。表 9-1 中的栏目 1 表示根据在第 8 章中建立的命名规则产生的法则代码; 栏目 2 表示在回溯测试期间, 法则在消除标准普尔 500 指数数据趋势中的平均

收益率; 栏目3和栏目4表示通过自举法和蒙特卡罗法(注意: 概率值并没有把数据挖掘偏差考虑进去)获得的单一法则的概率值; 栏目5~栏目7表示使用3种没有把数据挖掘偏差考虑进去的推理方法, 在0.05水平上的显著性。我们把任何一个具有显著性的法则用“*号”表示, 由于没有具有显著性的法则, 所以表格中就没有出现任何“*号”。栏目8和栏目9表示低于或者高于80%的置信区间。

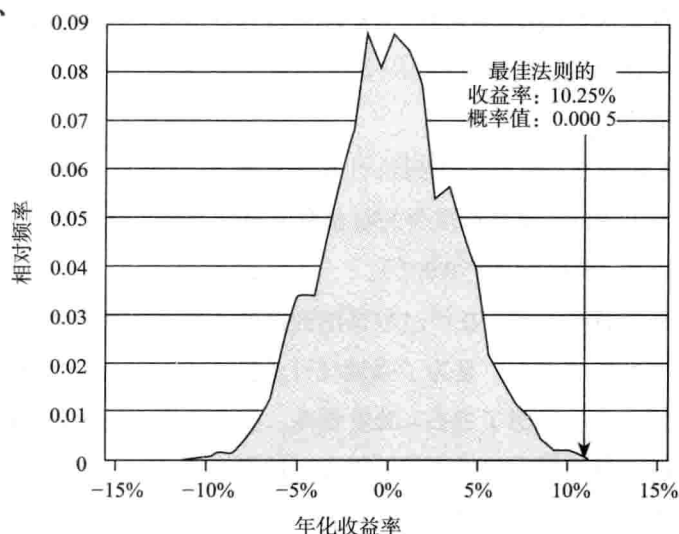


图 9-3 适用于对单一法则进行回溯测试的抽样分布

表 9-1 100 个拥有最佳绩效的法则

法则代码	收益率	自举法 p-value	蒙特卡罗法 p-value	B Sig	RC Sig	MC Sig	低于 80% 置信区间 CI	高于 80% 置信区间 CI
E-12-28-10-30	0.102 501	0.001 5	0.002				-0.029 32	0.235 27
D-8-4-10-60	0.102 395	0.001	0.001 5				-0.029 43	0.235 16
E-1-28-10-15	0.100 562	0.002 5	0.000 5				-0.031 26	0.233 33
E-12-21-20-15	0.098 838	0.000 5	0.001 5				-0.032 99	0.231 61
E-12-28-10-15	0.097 502	0.001	0.002 5				-0.034 32	0.230 27
E-12-24-20-15	0.095 904	0.000 5	0.001				-0.035 92	0.228 67
E-11-28-10-15	0.094 635	0.003	0.002				-0.037 19	0.227 4
E-11-39-10-60	0.092 282	0.001	0.001 5				-0.039 54	0.225 05
D-7-37-10-30	0.092 078	0.002 5	0.003				-0.039 75	0.224 85
D-7-37-10-15	0.090 196	0.004 5	0.004				-0.041 63	0.222 96

(续)

法则代码	收益率	自举法 p-value	蒙特卡罗法 p-value	B Sig	RC Sig	MC Sig	低于 80% 置信区间 CI	高于 80% 置信区间 CI
TT-37-5	0.089 024	0.004 5	0.003				-0.042 8	0.221 79
TI-20-41	0.088 62	0.004 5	0.004 5				-0.043 2	0.221 39
E-1-22-10-60	0.088 591	0.005 5	0.007				-0.043 23	0.221 36
TI-21-8	0.086 963	0.003 5	0.007 5				-0.044 86	0.219 73
0-10-28-10-30	0.086 939	0.004	0.002 5				-0.044 89	0.219 71
D-11-28-10-15	0.085 266	0.005 5	0.005 5				-0.046 56	0.218 04
D-7-38-10-30	0.085 206	0.005	0.006				-0.046 62	0.217 98
TI-28-12	0.084 912	0.003 5	0.003 5				-0.046 91	0.217 68
TI-24-8	0.084 021	0.007	0.005				-0.047 8	0.216 79
E-2-28-10-15	0.083 267	0.007 5	0.002 5				-0.048 56	0.216 04
E-5-36-10-30	0.083 087	0.005 5	0.006 5				-0.048 74	0.215 86
D-9-34-10-30	0.081 577	0.008	0.009				-0.050 25	0.214 35
E-11-15-20-60	0.080 58	0.007	0.007 5				-0.051 24	0.213 35
D-1-21-10-30	0.080 499	0.008	0.008				-0.051 32	0.213 27
D-8-32-10-15	0.080 48	0.006	0.008 5				-0.051 34	0.213 25
D-7-36-10-30	0.079 943	0.006 5	0.005 5				-0.051 88	0.212 71
D-8-27-10-30	0.079 601	0.007	0.008 5				-0.052 22	0.212 37
E-12-22-10-30	0.077 796	0.009	0.008				-0.054 03	0.210 57
D-1-24-10-30	0.077 582	0.01	0.012 5				-0.054 24	0.210 35
D-7-36-20-30	0.077 312	0.009 5	0.013 5				-0.054 51	0.210 08
D-8-33-20-15	0.077 101	0.009 5	0.008 5				-0.054 72	0.209 87
E-1-19-10-60	0.077 092	0.007 5	0.012 5				-0.054 73	0.209 86
D-8-32-20-15	0.076 434	0.010 5	0.009				-0.055 39	0.209 2
TI-8-18	0.076 144	0.009	0.01				-0.055 68	0.208 91
TI-26-137	0.075 767	0.009	0.010 5				-0.056 06	0.208 54
TI-24-5	0.075 732	0.008	0.009				-0.056 09	0.208 5
D-7-36-10-60	0.075 5	0.010 5	0.010 5				-0.056 2	0.208 27
E-1-28-20-15	0.075 41	0.013	0.013 5				-0.056 41	0.208 18
E-6-38-20-30	0.074 831	0.008 5	0.012 5				-0.056 99	0.207 6
E-12-24-10-15	0.074 809	0.011 5	0.013 5				-0.052 01	0.207 58
TT-37-3	0.074 756	0.012 5	0.01				-0.057 07	0.207 53
TT-38-5	0.074 579	0.013 5	0.012				-0.057 24	0.207 35
E-11-21-10-15	0.074 168	0.014 5	0.011				-0.057 66	0.206 94
D-9-4-10-60	0.074 09	0.012 5	0.012 5				-0.057 73	0.206 86
E-1-25-10-60	0.073 966	0.015 5	0.01				-0.057 86	0.206 74
E-1-18-20-60	0.073 439	0.015	0.015 5				-0.058 38	0.206 21
E-11-25-10-15	0.073 131	0.01	0.018				-0.058 69	0.205 9

(续)

法则代码	收益率	自举法 p-value	蒙特卡罗法 p-value	B Sig	RC Sig	MC Sig	低于 80% 置信区间 CI	高于 80% 置信区间 CI
D-10-32-20-15	0.072 548	0.020 5	0.014 5				-0.059 28	0.205 32
E-5-38-10-30	0.072 45	0.014	0.014				-0.059 37	0.205 22
D-11-28-10-30	0.072 45	0.018 5	0.013				-0.059 37	0.205 22
E-11-15-20-15	0.072 363	0.012	0.018				-0.059 46	0.205 13
E-2-24-20-15	0.072 286	0.018	0.011				-0.059 54	0.205 05
D-9-34-10-60	0.072 065	0.012	0.016				-0.059 76	0.204 83
D-7-11-10-15	0.071 859	0.013 5	0.011 5				-0.059 96	0.204 63
D-8-37-20-60	0.071 562	0.020 5	0.022				-0.060 26	0.204 33
E-2-21-10-15	0.071 546	0.017 5	0.015				-0.060 28	0.204 32
E-6-37-20-15	0.071 509	0.016	0.013 5				-0.060 31	0.204 28
TI-12-12	0.071 403	0.013 5	0.009 5				-0.060 42	0.204 17
E-12-21-10-15	0.071 14	0.016	0.019 5				-0.060 68	0.203 91
TI-23-5	0.071 013	0.015 5	0.011 5				-0.060 81	0.203 78
E-2-1-20-30	0.070 897	0.016	0.013				-0.060 93	0.203 67
TT-36-5	0.070 709	0.023 5	0.017				-0.061 12	0.203 48
E-6-38-20-15	0.070 681	0.017	0.021				-0.061 14	0.203 45
TI-21-12	0.070 672	0.017	0.010 5				-0.061 15	0.203 44
E-1-21-10-15	0.070 672	0.02	0.016				-0.061 15	0.203 44
D-8-38-20-30	0.070 308	0.02	0.022 5				-0.061 52	0.203 08
D-10-4-10-60	0.070 298	0.014 5	0.016				-0.061 53	0.203 07
TI-32-61	0.070 203	0.0215	0.013 5				-0.061 62	0.202 97
E-2-19-10-60	0.070 192	0.015	0.016				-0.061 63	0.202 96
TI-18-8	0.069 985	0.019	0.013 5				-0.061 84	0.202 75
TT-36-3	0.069 962	0.018 5	0.019 5				-0.061 86	0.202 73
D-8-23-20-60	0.069 884	0.017 5	0.019				-0.061 94	0.202 65
D-7-16-10-30	0.069 749	0.013	0.014 5				-0.062 07	0.202 52
D-7-34-10-60	0.069 721	0.016	0.023 5				-0.062 1	0.202 49
E-2-28-20-5	0.069 572	0.017	0.02				-0.062 25	0.202 34
D-9-37-10-30	0.069 551	0.023 5	0.019 5				-0.062 27	0.202 32
D-9-29-10-60	0.069 538	0.013	0.018 5				-0.062 29	0.202 31
D-9-36-20-30	0.069 506	0.022 5	0.018				-0.062 32	0.202 28
D-8-36-10-60	0.068 993	0.018	0.022				-0.062 83	0.201 76
D-8-33-10-15	0.068 802	0.025	0.023				-0.063 02	0.201 57
D-10-23-20-60	0.068 752	0.014 5	0.02				-0.063 07	0.201 52
D-7-37-20-30	0.068 628	0.02	0.021				-0.063 2	0.201 4
E-2-1-10-30	0.068 546	0.016	0.018				-0.063 28	0.201 32
D-8-31-10-30	0.068 31	0.019	0.023				-0.063 51	0.201 08
D-9-32-10-75	0.068 112	0.017 5	0.022				-0.063 71	0.200 88

(续)

法则代码	收益率	自举法 p-value	蒙特卡罗法 p-value	B Sig	RC Sig	MC Sig	低于 80% 置信区间 CI	高于 80% 置信区间 CI
E-5-20-10-30	0.068 035	0.018 5	0.011				-0.063 79	0.200 8
E-3-36-10-30	0.067 922	0.020 5	0.019				-0.063 9	0.200 69
TI-23-41	0.067 856	0.022 5	0.019				-0.063 97	0.200 63
E-2-9-10-15	0.067 833	0.018 5	0.021 5				-0.063 99	0.200 6
E-7-36-20-60	0.067 749	0.025 5	0.020 5				-0.064 08	0.200 52
TI-21-5	0.067 723	0.022	0.020 5				-0.064 1	0.200 49
E-6-38-10-30	0.067 453	0.022 5	0.024				-0.064 37	0.200 22
D-6-33-10-15	0.067 088	0.023	0.022				-0.064 74	0.199 86
D-8-34-20-60	0.066 988	0.025	0.023				-0.064 84	0.199 76
E-12-21-20-60	0.066 859	0.0175	0.025				-0.064 96	0.199 63
E-1-18-10-60	0.066 816	0.027	0.023 5				-0.065 01	0.199 59
D-5-15-10-30	0.066 604	0.024	0.017				-0.065 22	0.199 37

纳入数据挖掘偏差考虑的 3 种推理方法如下。

(1) B (boot) 是怀特真实检验 (WRC) 的改进版, 它吸收了由沃尔夫和罗马诺³ 提出的旨在改进测试力 (降低类型 II 错误的概率, 参见第 6 章) 的建议。

(2) RC (真实检验) 即 WRC 版本, 可以从 Quantmetrics 获得, 该版本并没有把沃尔夫和罗马诺的改进建议考虑在内。

(3) MC 是指由马斯特斯开发的蒙特卡罗置换法, 该方法把沃尔夫和罗马诺的改进建议考虑在内。

80% 置信区间的上限和下限是在数学理论的基础上得出的, 它由怀特⁴ 首先提出, 他用沃尔夫和罗马诺以及蒂姆·马斯特斯提出的算法把这一切转换成为计算机代码。置信区间充分考虑了数据挖掘偏差, 也就是说, 对于上限和下限在一起的联合置信区间在 0.80 的概率水平上, 包含了所有法则的真实收益率。这意味着, 如果我们对 1 000 个独立的数据集合重新进行案例研究的话, 只有一组历史市场数据包含了所有法则的真实收益率, 在对 6 402 个置信区间进行测试后, 只有 800 个可以包含所有法则的预期收益率。

我们也可以这样表达, 在这 1 000 个案例中, 有 950 个案例不会误导我们落入拒绝原假设的错误陷阱中, 而剩下的 50 个则有这个可能。碰巧在

为案例研究进行的特定实验中，没有一个法则显示出显著性。然而，通过把显著性设定在 0.05 的水平上，案例研究就会有 1/20 的概率出现类型 I 错误，即没有实际预测力的法则显示的概率值小于 0.05，从而导致对原假设进行了错误的排斥。

对案例研究的批评

优点

1. 在对法则进行评估时对基准的使用

案例研究的一个优点就是在评估法则的绩效时对基准的使用。正如我们在第 1 章中指出的那样，一个法则的回溯测试绩效只有和基准联系在一起时才有意义。绝对的绩效水平不包含任何信息。

案例研究使用的是合理的最低基准和不具有预测力的法则的绩效。根据这一标准，只有该法则运用统计显著性的优势击败了不具预测力的法则的绩效时，该法则才会被认为是有效的。显然，在特殊案例中还存在着更高的基准。假设有一种双重移动平均交叉法则被开发出来，那么其合理的基准应当是法则的传统版本所产生的绩效。

2. 对市场趋势有效性的控制

我们在第 1 章中曾经指出，案例研究利用消除趋势的市场数据计算出来的法则收益率来消除失真的绩效。正如我们曾解释的那样，当持有多头头寸/空头头寸偏差的法则获得通过在回溯测试期间具有净上升和净下跌趋势的市场数据计算出的收益率时，失真才会出现。

3. 对前视偏差的控制

法则研究假设信息的可得性在市场头寸已经进场或者已经离场时并不是真正可以获得的，它们都会受到前视偏差的影响，我们把这种现象称为“对未来信息的泄漏”。如果需要用收盘价的信息计算法则的信号，那么把收盘价视为进场或者离场的假设就是不合理的。在进场或者离场时能够被

合理假设的第一个价格就是下一个价格。如果用于案例研究的数据是每日的频率,那么第一个合法的执行价格会一直到发出信号的第二天市场开盘时才会出现。

如果法则使用的数据序列具有滞后性,如共同基金的现金统计数字,或者像政府经济统计数据这种经过修正的数据,都会产生前视偏差。如果事情真是这样的话,那么使用的滞后值就必须把滞后或者修正考虑进去。

案例研究通过假设发出头寸反转信号的第二天开盘时的进场和离场情况的方式来避免前视偏差。另外,没有数据序列受到滞后或者修正的影响。

4. 对数据挖掘偏差的控制

我们对流行的技术分析中的一些法则进行了各种类型的显著性测试。但是,这些测试并没有对由普通抽样的错误而导致法则产生利润的可能性做出说明。这是非常严重的疏漏,但是我们可以通过应用普通假设检验的方法来修正。

然而,在法则已经进行过回溯测试的情况下,普通的显著性测试还是比较适合的。而当我们对很多法则进行回溯测试后并把最佳法则挑选出来时,普通的假设检验就会让那个最佳法则比其实际产生(对原假设的错误拒绝)的统计显著性更加突出。为了避免类型 I 的错误,我们就要应用诸如 WRC 和蒙特卡罗置换法等更加先进的假设检验法,这两种方法我们在第 6 章中做过介绍,本书的案例研究使用的就是这两种方法。正如我们之前指出的那样, WRC 和蒙特卡罗置换法显示,在 6 402 个经过测试的法则中,那个产生最佳绩效的法则并没有显示出统计显著性。

5. 对数据探测法的控制

也许我们用“事先研究探测法”这个名称来表示数据探测法偏差会更加恰当,当数据挖掘器使用之前的研究结果来选择对哪个法则进行测试时,事先研究探测法偏差就会出现。换句话说,他们使用的是在这之前已经被证明是成功的法则。由于进行测试的法则的数量对于确定数据挖掘偏差程度来说是非常重要的因素,所以对包括在新的数据挖掘风险中的法则的实际统计显著性是无法做出评估的,这是因为该法则已经在之前的研究中取

得了成功。

案例研究通过明确包含任何由其他研究者发现的法则来减少事先研究探测偏差。很多用案例研究中的 6 402 个法则与由其他研究者进行研究的法则相似，而且他们之间的相似点却是一致的。就像我们在第 8 章中解释的那样，我们通过对所有可能出现在一组特定的参数值内的所有法则进行组合计数的方法来实现对法则集合测试的目的。例如，趋势法则以可能出现的 11 个通道突破追溯参数值和 39 个可能的时间序列为基础。所有 429 (11×39) 种可能出现的法则都包括在案例研究中。如果这些法则当中的任何一个法则碰巧就是由之前一些研究发现的法则，那么案例研究的内容就纯属偶然，而不是分析的结果。因此，对所有的法则进行解释就可以导致最佳绩效法则的出现。

负面属性

1. 不考虑复杂法则

也许案例研究的最大缺陷就是没有把复杂法则考虑在内。复杂法则是通过把存在于单一法则中的大量信息进行合并与压缩而形成的。复杂法则具有几个致命的缺陷。

首先，有些技术分析的实践者（不论是主观的还是客观的）在做出决策的时候只局限于一种法则。从这种层面上来讲，案例研究并没有把主观分析师和客观分析师实际的操作进行重叠。因此，主观分析师在他们对于多重法则结构的信息模式的解读能力方面受到了严重约束。

其次，复杂法则在对诸如金融市场这样复杂的问题做出预测方面，往往能够产生最佳绩效。对单一法则的非线性合并使得复杂法则比从独立要素中总结出来的信息包含更多的信息，这就能够让法则遵循能够保证问题和解决方法必须具备相同复杂性的阿什比必要多样性定律（Ashby's Law of Requisite Variety）。⁵ 关于预测，则暗示我们旨在试图对复杂系统（如金融市场）的行为进行预测的模型也必须是复杂的。总的来说，复杂法则的优势甚至可以是线性模式，它是通过对包括单一和复杂法则在内的 39 832 个

法则进行分析⁶时被证明的。这组法则集合对道琼斯、标准普尔 500 指数、纳斯达克以及罗素 2000 指数这 4 种股票指数从 1990 ~ 2002 年的表现进行了测试。在全部的法则集合中, 3 180 个(或者 8%) 法则是复杂法则。在产生统计显著性利润的 229 个法则当中, 有 188 个(或者 82%) 法则是复杂法则。研究工作把数据挖掘偏差考虑了进来, 并利用 WRC 方法计算统计显著性水平。顺便说一句, 研究显示没有一个产生统计显著性的法则可以超过道琼斯工业指数或者标准普尔 500 指数。

案例研究仅限于使用单一法则, 以方便对其范围的管理。然而, 我却相信 在 6 402 个法则中, 有一些法则是可以证明其显著性的, 当然, 我对此是过于自信了。

2. 只考虑多头 / 空头反转法则

为便于管理案例研究, 法则集合仅限于一种类型的二进制法则: 多头 / 空头反转。正如我们在第 1 章中讨论的那样, 这种限制将会扭曲一种法则打算表达的技术分析概念。从整体上来说不仅是对被普遍认可的二进制法则的限制, 而且是对多头 / 空头反转法则的限制。对于市场永远处于无效状态和持续发出盈利机会这种不可能的假设做出的反转法则, 通常都要持有市场头寸的要求似乎是合理的。反过来, 当股市投资比率被证明与更多的合理假设(市场只是偶尔处于无效状态)相一致的时候, 我们在法则的识别上会更具有选择性。因此, 三态法则可以存在多头 / 空头 / 中性头寸, 或者是只包含多头 / 中性和空头 / 中性头寸的二进制法则会表现得更好, 但是这种类型的法则没有被应用到案例研究中。

3. 对标准普尔 500 指数交易进行限制的法则

同样出于对案例研究范围进行限制的目的, 这些法则只被应用到标准普尔 500 指数。大规模的法则研究⁷指的是, 早期没有发现任何对标准普尔 500 指数产生影响的单一或者复杂的法则。当选择标准普尔 500 指数作为交易市场时, 我没有意识到进行更大的法则研究, 因为这样做可以为纳斯达克和罗素 2000 指数发现有用的法则。一方面, 如果我意识到进行更大的法则研究, 并且把纳斯达克和罗素 2000 指数作为交易市场, 我就会对事

先探测研究产生愧疚感，并且向从案例研究中获得的概率值的真实性做出让步了。另一方面，要是我在不知道案例研究可以发现优秀法则的情况下选择了纳斯达克和罗素 2000 指数，那么这些法则的发现就是顺理成章的，且概率值也是精确的。

案例研究扩展的可能性

前面提到的各种类型的限制为案例研究提供了几种扩展的可能性。

案例研究对不成熟市场的应用

Hsu 和 Kuan 的研究指的是，早期在纳斯达克和罗素 2000 指数中发现的具有统计显著性的法则，技术分析法则在罗素 2000 指数这种尚不成熟的股票市场中往往会更加成功。因此，对案例研究进行的其中一种扩展就是，把案例研究使用的法则集合应用到证券发展不成熟的国家的股票市场指数中去。如果想把原始的数据序列设计成为上涨 / 下跌统计数字、上涨 / 下跌量、新高点 / 低点等可以让人接受的指标的话，这种做法还是非常实际的。一部文献资料显示，相对于发达的市场和更加有效的市场而言，新兴的股票市场也许会更适合预测，而大部分研究⁸认为这种观点是非常合理的，新兴股票市场的现实局限性在于其短暂的市场历史以及较少的历史观察资料结果。

改进的指标和具体法则

法则的盈利性取决于其所依靠的指标的信息含量。一种指标越能有效地把信息含量大的市场行为特点进行量化，该法则的潜在盈利性就越高。用于案例研究的 6 402 个法则包括 3 个技术分析主题：趋势、极端情况以及背离。指标曾经被认为是量化这些主题的合理方法，毫无疑问，它们还具有改良的余地。正如我们本章后面内容指出的那样，技术分析的实践者未来最重要的任务将是以提供更多信息的方式来发展能够对市场行为进行量化的指标。

例如，我们在第8章中提到的，背离有可能通过基于协整关系的指标进行更好的衡量。⁹如果他们使用埃勒斯¹⁰论述的费希尔转换，那些旨在衡量极端情况的指标也许就会被证明是有用的。他指出，基于通道规范化指标的频率分布具有不符合需要的特性，具体来说，就是极端值可能大大超过中间值的指标范围。这对于缺少旨在发现罕见的极端指标的我们来说，几乎是不可能完成的任务。然而，通过把费希尔转换应用到通道规范化数据当中，埃勒斯声称合成的指标会更接近于正常分布。一个拥有正常分布的指标具有其极端值非常罕见的特征，这也许会使它们具备更多的信息。由于埃勒斯并没有提供费希尔转换有效性的数据支付实证，所以它只是一个值得进一步研究的有趣猜想。

把复杂法则考虑在内

复杂法则是指为了论证预测力而从两个或者更多的单一法则合并之后得出的信息中推导出来的。在推导复杂法则的过程中，出现了两个主要问题：我们需要合并哪些法则以及如何对它们进行合并。这里的“如何”是指用于将单一法则的输出经过合并而产生出复杂法则输出的数学或者逻辑形式。

最简单的方法就是线性组合。线性组合是以把简单法则的输出进行累计为基础的。它处理每一个简单法则的输出就好像它对复杂法则的输出做出了独立的贡献一样，也就是说，每一个单一法则的贡献没有受到其他法则输出的影响。最简单的线性组合对每个因素都做了同样的加权处理。因此，当两个二进制反转法则都持有空头头寸的时候，两个二进制反转法则的简单线性合并可以假设为-2的数值；当两个法则分别持有多头头寸和空头头寸的时候，该数值为0；当两个法则都持有多头头寸的时候，该数值则为+2。更为复杂的是，尽管线性合并没有一定的优势，但是它可以根据自己的预测力对每一个法则指定不同的加权数。每一个法则的相对加权数通过系数 a 表示。根据最小预测错误或者财务业绩等品质因数，加权集合可以用优化的方式发现。经过加权的线性合并的形式请参见下面的公式：

线性组合是加法

$$Y = a_0 + a_1 r_1 + a_2 r_2 + \cdots + a_n r_n$$

式中, Y 代表复杂的线性法则的输出; r_i 表示法则 i^{th} 的输出; a_i 表示法则 i^{th} 的加权数; a_0 表示常量——截距。

与线性组合相反, 非线性组合考虑的是法则之间的截距。换句话说, 经过合并的价值 Y , 并不是单一法则数值的加权和。在非线性合并中, 任何单一法则对合并的输出值 Y 的基值在某种程度上取决于其他法则的输出值。在线性组合中, 每个单一法则对 Y 的贡献是不受其他法则的数值的影响的(独立的)。非线性组合的优势是, 它可以表达更加复杂的信息模式。然而, 它也是存在一些缺点的, 即非线性组合也许不是一个能够清楚表达出来的方程。由于我们对法则的输出是如何进行合并的(相互作用)还不知晓, 因此非线性模型有时是指黑盒子公式(black-box formulas)。由神经网络软件产生的非线性组合将会作为黑盒子公式的实例。

线性和非线性组合之间的区别, 在我们只考虑连续变量而不是二进制输出时, 是很容易描绘的。连续变量可以假设任何变量处于某一范围内, 而不是把离散值限制在范围为 $(+1, -1)$ 这样的二进制法则中。当线性模型用图像表示的时候, 如图 9-4 所示, 我们假设代表复杂法则输出 Y 的响应面是平面形状, 它的表面是光滑的, 没有山丘和峡谷。另一种说法是, 每一个法则的输入值 (X_1 和 X_2) 响应面的倾斜度, 在整个输入范围内保持不变。我们把输入 X_1 的响应面倾斜度用系数 a_1 表示, 而 X_2 的响应面倾斜度用系数 a_2 表示。 a_1 和 a_2 一般情况下是从一种被称为回归分析的统计方法中的一组数据中被发现的。把 a_1 和 a_2 与第三种系数 a_0 进行合并组成截距 Y , 让平面能够更好地穿过用于评估模型系数的数据点集合。响应面的倾斜度是常量(平面)这一事实告诉我们, 在两个输入量之间是不存在相互作用的。换一种说法就是, 输入量 X_1 和 X_2 是以相加的方式进行相互作用的。

在非线性模型中, 两个输入量之间非叠加性的相互作用是被允许的, 这就意味着数值 Y 并不是单一法则数值的加权和。在非线性模型中, 任何输入量(指标)的基值是以其他输入量的数值为基础的, 这就证明了具有

山丘和峡谷的响应面。换句话说,对于可以是任何单一输入量 X_1 的模型响应面的倾斜度来说,其范围都不是平面的。存在于覆盖输入量范围倾斜度中的变量可能会产生出如图 9-5 所示的响应面。

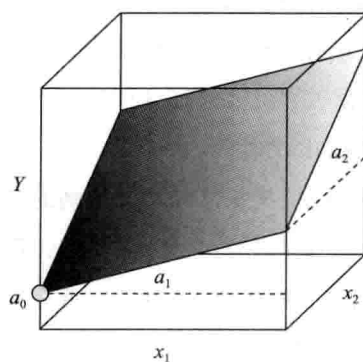


图 9-4 $Y = a_0 + a_1x_1 + a_2x_2$

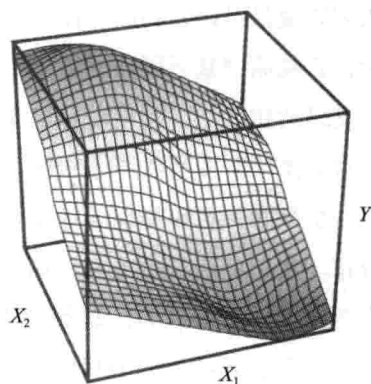


图 9-5 $Y = f(X_1, X_2)$

1. 在一个主题之内的线性组合

尽管通过加法进行法则的合并非常简单,但是通过线性方式形成的复杂法则却更加有用。Hsu 和 Kuan¹¹ 的研究试验了 3 种建立复杂法则的方法。其中的两种方法以线性形式为基础:投票法则和部分头寸法则。第三种是学习法则,它的目的不是将法则进行合并,而是从一组相同的法则集合中挑选出最近产生最佳绩效的单一法则。基于投票法则和部分头寸法则的复杂法则,实际上就是所有在特定主题内的简单法则的叠加组合。这些主题的实例就是以通道突破和能量潮的移动平均为基础的所有法则。基于投票方案的复杂法则的输出值是通过在一个主题内的(通道突破)所有法则进行投票,并且提取一个能够反映大多数投票的单一头寸推导出来的。例如,如果大多数法则显示出了空头头寸,那么单一的空头头寸就会被接受。基于部分头寸概念的复杂法则的输出值是通过在某一主题内的所有单一法则进行投票的基础上得出的。然而,被采用的头寸的规模应当与多头/空头法则(多头-空头/全部法则的数量)的净百分比相等,如 Hsu 和 Kuan 研究的 2 040 个能量潮(OBV)法则。如果其中的 1 158 个法则发出空头信号,而 882 个法则发出了多头信号,那么数值为 0.135 的单一空头

头寸就会被采取 $(882-1\ 158)/2\ 040=-0.135$ 。

因此,一种符合案例研究逻辑的延伸就会产生以投票法则和部分头寸计划为基础的线性复杂法则。这对于在给定主题范围内(如背离法则)的法则来说是可以做到的。在技术分析领域被普遍使用的一种指标是扩散指标,它的原理基于同样的概念。最广为人知的一种扩散指标就是,纽约股票交易所位于200日移动平均线之上的股票的百分比。在这种情况下,这个百分比从本质上来讲是一个没有经过加权的线性组合。

扩散指标可以被简单地认为是具有多头/空头头寸的背离法则的净百分比。我们使用投票方案,如果超过50%的背离指标处于+1状态,那么多头头寸就是有效的;如果超过50%的背离指标处于-1状态,那么空头头寸就是有效的。如果利用部分头寸法则的话,那么部分多头/空头头寸就是有效的,而且它们会随着净多头/空头头寸百分比的变化而进行调整。需要特别注意的是,从案例研究的6 402个法则当中产生的扩散指标需要消除相反的法则,因为它们得出的结果通常会假设一个为0的数值。因此,以三类趋势法则为基础的扩散指标只能通过TT法则建立,而TI法则将会被排除。同理,以极端情况和趋势转换为基础的扩散指标则只能以类型1、2、3、4、5、6为基础,而类型7~类型12则必须被排除,因为它们是类型1~类型6的反转形态。

2. 通过机器建立的非线性组合

通过把单一法则进行合并或者以非线性或非叠加的方式形成的复杂法则也可以从案例研究的法则当中获得。然而,这就需要使用自制的机器学习(数据挖掘)系统,如神经网络、¹²决策树和决策树总体、多元回归曲线、¹³核回归、¹⁴多项式网络、¹⁵遗传算法、¹⁶支持向量机¹⁷等。这些系统采用归纳推广,通过对大量历史样本的研究,对复杂的非线性法则(模型)进行合成。每一个样本都是一个案例,并且以截止到某一特定日期的一组指标值为特征,模型以结果变量的数值进行预测。

复杂的非线性法则可以通过两种方式推导:一种是以二进制形式提出的法则作为备选的输入量,另一种方法就是提出法则所使用的指标。

就我所掌握的知识来说,目前还没有可得的结果用于产生复杂的非线性

性法则，而该法则是对数据挖掘偏差问题进行稳健的显著性测试的直接体现。在缺少保护的条件下，数据挖掘器对复杂的非线性法则的寻找必须依靠由3个数据集（培养、测试和评估）组成的复杂样本外测试形式所提供的保护来完成。这一点我们将在后面的章节中讨论。

与其寻找对法则复杂性进行优化的方法，倒不如限制其发现一个复杂性（上涨和下跌）被固定的法则的最优参数值。上涨趋势可以让复杂性优化法增加发现最佳法则的机会，而增加寻找范围的下跌趋势则可以大幅增加在最佳法则被挑选出来之前所考虑的法则数量，这就增加了具有最佳绩效的法则出现过度拟合的概率。这就是说，一个法则不仅表达了样本数据的有效性形态，而且还描述了其随机效应和随机噪声。过度拟合的法则并没有被很好地概括。也就是说，它们的样本外绩效相对于样本内绩效表现较差。

为了防止在复杂性优化法被允许的情况下出现过度拟合的风险，我们通常使用3个数据段：培养、测试和验证。回想一下我们在第6章提到的内容，当数据挖掘仅局限于发现一个复杂性被固定的优化¹⁸参数值的时候，我们只需要两个数据段：培养和测试。

在寻找最佳复杂性的非线性法则的过程中，机器学习算法通过培养集和测试集在两个不同的圈子中进行循环。内圈是在给定的复杂性水平下寻找最优参数，而外圈则是寻找复杂性的最优水平。一旦发现了最佳的复杂性法则，我们就在被称为第三个数据段的“验证集”中对其进行评估。这个数据将会一直被保存到参数值和最佳绩效法则的复杂程度被发现的时候。由于在发现的过程中并没有使用“验证集”，则在这个数据集中最佳法则的绩效就是其未来绩效的无偏差估计（见图9-6）。

上述的三段计划可以实施到前行偏差当中。当在第一组数据折叠（培养、测试和验证）中发现最佳法则之后，分成3个部分的数据窗口向前滑行，并把整个过程在下一个折叠中进行重复（见图9-7）。

分成3个部分数据窗口的移动证明了一些额外的解释。市场行为被假设是一组系统性行为（周期性的模式）和随机噪声的组合，通过增加法则的复杂性提供符合给定数据段的法则是很有可能的。换句话说，给定足够

的复杂性, 就可以让该法则在每一个市场的最低点买入或者在最高点卖出的。这就是一个糟糕的想法。¹⁹ 在过去的数据中出现的完美时机仅仅会得到受到噪声影响的法则这一结果。换言之, 完美的信号或者任何接近它们的信号几乎肯定意味着该法则受到干扰的程度, 且还是对过去的随机行为进行(过度拟合)的描述。

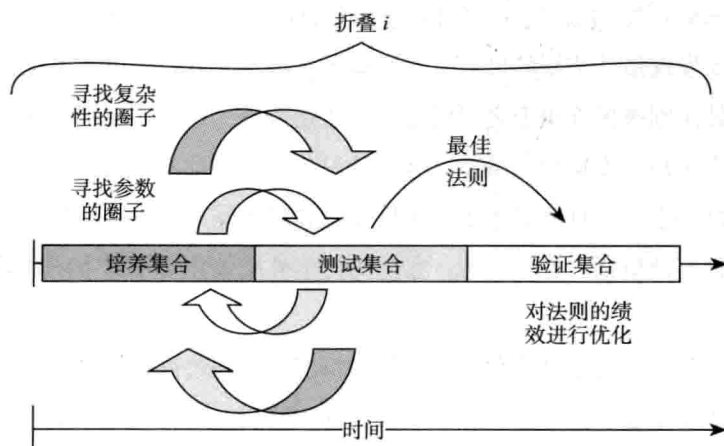


图 9-6 寻找最优参数和最佳法则的复杂性

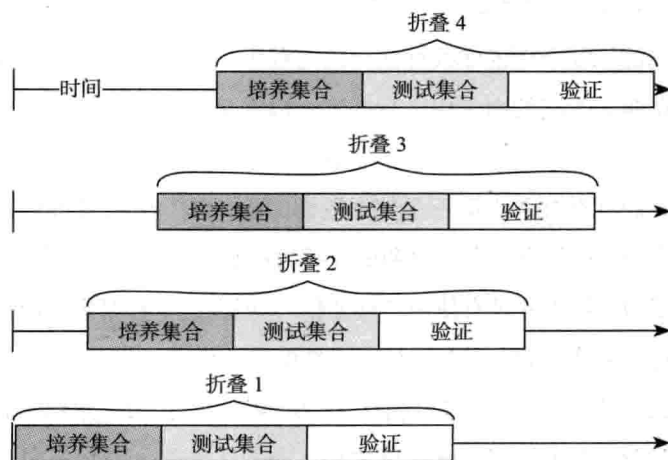


图 9-7 通过向前移动寻找复杂性

当把法则应用到对数据段的测试中时, 过度拟合现象就会很明显。所以, 其绩效也要比在培养数据中的绩效差, 这是因为在培养集中发现的合理形态会在测试集中重新出现, 但是在培养集中的噪声就不会再出现。我

们可以由此推测出不会在测试集合中再次出现的培养集合的盈利性，是最有可能出现的过度拟合的结果。

在对复杂性进行优化的过程当中，当法则的绩效达到测试集的顶峰时，法则出现过度拟合的信号就会产生，然后开始逐渐降低。对复杂性进行优化的过程如下：我们以一个复杂性较低的法则开始，并且使用不同的参数值将其在培养数据集中反复运行，每次运行会产生一个绩效值（如收益率）。我们把在培养数据中产生最高绩效的参数称为“最佳法则 1”，然后再将其在测试数据集中运行，并且记录下它产生的绩效。至此，我们就完成了第一个复杂性循环。我们在进行第二个复杂性循环时使用的法则，是一个比最佳法则 1 等级要高的法则。也就是说，它拥有更多的参数。我们再重复上面的步骤，得出这个更加复杂的法则在训练集中的最佳参数值。我们把处于这个层面的最佳法则称为“最佳法则 2”，最佳法则 2 之所以会在培养集合中获得比最佳法则 1 更好的绩效，是因为其具有额外的复杂性。我们然后把最佳法则 2 的绩效通过测试集进行衡量。我们利用采自培养集的复杂性不断增加的法则继续这个测试过程，然后再将它们进行测试集进行评估。一般说来，随着法则复杂性的增加，绩效的表现也会随之改变。这种绩效的改变就是在培养集中更加复杂且有效的模式被发现的迹象。然而，从某种角度来说，随着法则复杂性的增加，其在测试集中绩效将开始下降。这种情况表明对复杂性的最新增量可以对培养集中的随机效应（一次性噪声）进行描述。换言之，该法则现在已经出现了过度拟合。如果对法则复杂性的增加能够在这个点之前停止的话，那么法则就无法捕捉到所有数据的有效形态。我们把这种法则称为不充分拟合。图 9-8 显示了过度拟合和不充分拟合的情况。请注意，复杂性的增加通常会在培养集中产生经过

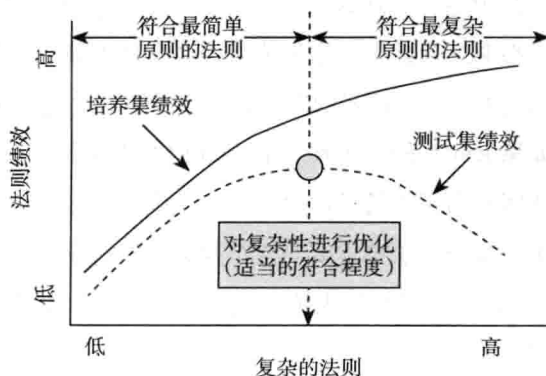


图 9-8 对复杂的法则进行优化

改进的绩效, 这种情况是通过在培养集中持续增加的绩效显示的。我们还要注意到过度拟合和不充分拟合之间的分界线是由测试集绩效的顶点并且开始下降的那个点来决定的。

复杂性概念是抽象的, 对有些读者来说也许是新鲜事物。为了对其进行说明, 我们需要考虑下面的人工指导, 而不是考虑由机器指导的对复杂性的优化。我们首先从基本的两条移动平均线交叉的反转法则开始。²⁰ 该法则由两个参数确定, 即短期移动平均追溯范围和长期移动平均追溯范围组成。首先, 我们把这两组移动平均所有可能出现的参数组合在培养集中运行。如果每条移动平均能够允许 10 个参数值, 那么就可以进行 100 次的强力搜索。²¹ 对培养集中的最高绩效(最佳法则)进行参数合并, 然后在测试数据集中运行。假设最佳法则在培养数据中可以获得 30% 的年收益率, 而在测试数据中获得 15% 的收益率。到目前为止, 我们只有最普通的参数优化, 这是因为测试数据一直以来只应用了一次, 15% 的收益率就是法则未来绩效的无偏差估计。

现在, 假设我们对 15% 的年收益率不满意, 我们就要对更广泛、建立在简单的两条移动平均线交叉的法则之上的复杂法则集合进行数据挖掘。²² 我们现在回到该计划的初级阶段, 并且提出一个通过在相对强弱指标(衡量价格变动率的指标, RSI)的基础上增加一个法则的复杂法则。在这个样本中, RSI 法则的增加, 使得原始的两条移动平均线交叉的法则由一个从多头到空头或者从空头到多头的反转法则向一个可以在所有市场之外持续的法则转变。达成这一目的的方法之一就是把 RSI 指标视为一个过滤器: 如果 RSI 指标大于临界值, 即 RSI 最大值, 也就是由双重移动平均线交叉法则发出卖出信号的时候, 一个中性头寸就会被用来取代空头头寸; 相反, 如果 RSI 指标小于临界值, 即 RSI 最小值, 也就是由双重移动平均线交叉法则发出买入信号的时候, 一个中性头寸就会被用来取代多头头寸。这种更为复杂的法则需要对 4 个参数进行优化: 短期移动平均范围、长期移动平均范围、RSI 最大值以及 RSI 最小值。大量的参数合并将会与培养数据进行竞争。

正如我们期待的那样, 复杂性的增加让包含 4 个参数在内的法则在培

养数据中获得了更高的收益率。假设使用在培养数据中产生的 50% 的收益率，与之前讨论过的只有两个参数组成的法则产生的 30% 的收益率进行比较，在培养数据中增加的绩效是对增加复杂性（增加参数）的可预测结果。然而，我们惊喜地发现，包含 4 个参数在内的法则在测试集中的表现同样优秀，它与只有两个参数组成的法则产生的收益率分别是 20% 和 15%。当更加复杂的法则在测试集中有更好的表现时，我们有理由假设增加的复杂性可以成功地开发出市场行为的额外系统性效应。

因此，我们再次回到该计划的初级阶段，用没有增加复杂性的方法来提高法则的绩效。这次我们加入了新的参数。这种方法可以被视为另外一个过滤器，或者是可以选择的进场信号。随着复杂性的每一次增加，参数值所有可能出现的组合形态都会在培养数据中进行评估，并且把产生最佳绩效的参数组合与测试数据进行竞争。增加复杂性的过程直到在测试集中的绩效开始下降时为止。过度拟合和不充分拟合之间的分界线进行交叉的地方就是这个信号出现的地方，增加复杂性已经开始与培养数据中的噪声相一致了。最佳法则就是在测试集中绩效最佳的那个拥有最高复杂性的法则。

那么第三个数据段——验证集的目的又是什么呢？到目前为止，我们还没有触及这一层面。聪明的读者可能已经意识到，通过对在测试集中的复杂法则进行优化获得的绩效肯定会出现偏差。通过对测试集的再度观察，我们发现对复杂性进行优化的法则降低了正偏差，即数据挖掘偏差。因此，为了在将来获得经过优化的复杂法则的预期绩效的无偏差估计，我们必须进入验证集。因为在对优化参数值或者对复杂性进行优化的寻找过程中，并没有使用验证集，它的作用有待进一步验证。

技术分析的未來

技术分析当前正处于十字路口，它的未来取决于实践者所选择的道路。一种是传统的技术分析方法，它建立在不可检验的命题、轶事证据以及主观主义者所持有的传统的天真的理论直觉分析的基础之上。而另外一种则

是被我称为“实证技术分析 (EBTA)”的科学方法，这种方法只局限于那些可以进行检验的方法、客观的实证以及用适当的统计学推断工具进行严格且持有怀疑态度的分析。因此，实证技术分析在愚蠢的轻信与不断的怀疑之间开辟了一条道路。

我现在做出了一个预测：如果不能够把技术分析现代化，那么其就有可能被边缘化。历史上有充足的例子可以证实这个观点：占星术由于推动了天文科学的诞生而处于被边缘化的状态，而天文科学的日益繁荣则成为了当今世界的主流学科。

如果技术分析还在坚持其传统的非科学道路的话，那么是把它和其他古老的习俗一起抛弃，还是因为它采用了实证技术分析而可以保留一些重要的和相关的内容呢？历史上的先例表明这些不一定对我们有利。科学的历史指出，一些曾经致力于某种模式的实践者最终选择了放弃。这一点与我们在第2章提到的实证是一致的，即有些已经建立的观念具有极端的适应性。研究显示，一种观点可以承受完整的证据的质疑，而这些证据又是最初形成这种观点的。如果技术分析想要维持其传统的道路，那么这将会是非常遗憾的。我相信，这将会不可避免地导致技术分析放弃其最初支持过的严谨的学科，比如说经验主义和行为金融学。

当然，这也并不是说技术分析就会失去其所有的支持者。那些预言者和告知者肯定还会拥有不少的客户，毕竟对未来的预测是世界上排名第二的最古老的职业。²³有些听众之所以对此过于自信，是因为他们需要一种来自心灵的慰藉来减少由于对未来的不确定性而引起的焦虑。根据斯科特·阿姆斯特朗这位理论预测专家的观点，他认为对预言服务的强烈需求不一定非得归结为对预测准确性的要求。在他的论文 *The Seer-Sucker Theory: The Value of Experts in Forecasting*²⁴ 中，阿姆斯特朗辩论道，即使专家们的预言缺乏准确性，但是人们仍然看重专家的观点。消费者实际的购买行为是由于其在不确定性当中虚幻的减少以及把责任全部推诿给专家的结果。阿姆斯特朗之所以把这些消费者视为容易上当受骗的傻瓜，是因为他在论文中的一连串引用表明，专家的预测比那些不是专家的人做出的预测好不了多少。这种对预测准确性的缺失还出现在诸如金融预测等学科

当中。

一篇由阿姆斯特朗引用的开创性研究为阿尔弗雷德·考勒斯 (Alfred Cowles)、²⁵ 汉密尔顿 (Hamilton) 所援引, 阿尔弗雷德·考勒斯对 20 世纪 30 年代早期的市场专家的业绩记录进行了研究, 而汉密尔顿是《华尔街日报》的主编。考勒斯的研究表明, 1902 ~ 1929 年, 尽管汉密尔顿在其读者中享有盛誉, 但是市场专家对方向变化的预测有 50% 是错误的。考勒斯还对 20 家保险公司、16 家投资咨询所以以及 24 家金融出版物的预测绩效进行了评估, 但是这些机构并没有留下引人注目的业绩记录。同样的结论还可以在由《赫伯特金融摘要》提供的统计资料中得出, 该摘要目前提供了由时事通讯社推荐的超过 500 种投资证券组合的绩效。在赫伯特进行的一项研究中, 他对 57 个时事通讯社从 1987 年 8 月到 1988 年 8 月这 10 年间的表现进行了跟踪。在这段时间里, 有不到 10% 的时事通讯社跑赢了威塞尔 5000 指数的复合收益率。

阿姆斯特朗还声称, 有些级别较高的专家增加了预测的准确性。因此, 那些基于最便宜的预测买入而大赚一笔的客户就可以与最值钱的专家们相提并论了。那些成本最便宜的预测的准确度与成本最贵的预测, 或者是通过适度的投资而获得的准确性是一样的。

最近有迹象表明, 华尔街的那些精明的消费者不再愿意为传统的技术分析而付出代价。2005 年, 华尔街两家最大的经纪商关闭了传统的技术分析研究部门。²⁶ 这到底是暂时现象还是一种趋势的开始, 还有待观察。

实证技术分析：科学途径

采用科学的方法可以让技术分析马上从中受益匪浅。我们可以通过把主观的技术分析重新用可以检验的客观方法表示, 或者干脆把两者全部抛弃的方法使主观的技术分析被淘汰, 从而把技术分析转换成为合理的科学的行为。只有能够接受检验的思想才有希望, 而那些能够用客观实证证明自己的技术分析才可以加入到其知识体系中来。

但这并不是说当实证技术分析被接受的时候, 所有的答案就都一目了然, 且所有的争论就都停止了——事实远非如此。争论只是科学过程中的

一个重要环节。即使在自然科学领域中，重要的问题正在寻求解决，把可检验性剥离之后做出的有意义假设。

实证技术分析是有局限性的。其中之一就是它不具备实施可控试验的能力，这一点是科学研究的黄金标准。在可控试验中，可以通过保持其他变量不变的方法，把一个变量的效应隔离并研究。技术分析将不可避免地成为一种以历史为基础产生新知识的观察科学。

就这一点而言，技术分析并不是单独存在的，考古学、古生物学以及地质学同样依靠于历史数据。然而，它们之所以被称为学科，是因为这些学科涉及可检验的问题，并且通过对历史客观实证的观察将无用的思想排除出去。

统计程序在一定程度上改善了试验控制的缺失问题。例如，当福斯贝克²⁷对共同基金经理的乐观与悲观进行量化时，他利用回归分析的方法把短期利率对准备金产生的混淆效应排除了出去。通过消除短期利率对共同基金准备金产生的强大影响，他可以获得对基金经理进行不受任何干扰的心理测量，我们将其称为福斯贝克指数。运用类似的方法，雅各布斯和利维²⁸使用多重回归分析将多变量的混淆效应进行最小化，这使他们可以得到对相关股票的绩效进行预测所需的经过净化的指标。例如，他们可以衡量纯收益率反转效应，²⁹以此来表明他们是最给力且不矛盾的预测者之一。换句话说，他们发现了起作用的技术分析，也就是说，在过去的几个月中表现疲弱的股票，有可能在未来的几个月中获得相对不错的绩效。

无法控制的检验的另外一个缺陷是它不能产生新的数据样本。技术分析的研究者拥有一套可以重复使用的完整的市场历史资料，这将会导致数据挖掘偏差问题。在给定的数据集中进行测试的法则越多，偶然出现法则与数据相符合的概率就越大。我们可以把不能产生新的观察结果集合的情况，通过实验科学来实现。实验科学强调了统计学推断方法的重要性，因为它把数据挖掘偏差问题也考虑了进去。

人机交互中的专业法则

我相信技术分析的未來一定会通过技术分析的专家和计算机为我们提

供最好的服务。这个观点算不是我新近提出的，30年前我在费尔森³⁰的著作中就提出了这个观点，并且让技术分析的实践者在同时利用先进的数据模型和机器学习优势方面处于前沿地位。³¹在过去的10年中，在同行评审期刊中出现的大量学术论文显示了把先进的数据挖掘方法应用到技术指标和基本面指标中的重要性。我们已经找到了有效的客观形态。³²

这种交互利用的是人类专家和计算机之间的协同作用。这种协同作用之所以会存在，是因为这两个实体处理信息的能力是值得称赞的，因为计算机所擅长的部分正是人类最薄弱的地方，反之亦然。这里指的是在特殊领域中装备数据挖掘软件的计算机和拥有专业知识的人类，如技术分析、细胞学、油田地质学等领域。

也许我把人类比计算机更具有创新能力这个问题过于简单化了。人们可以提出问题，并且通过假设的解释把不同的事实组成一种形态。然而，人类所具有的这种独特能力也有不好的一面：它让我们受到蒙蔽。人类之所以会上当受骗，是因为人类的思想已经进化到了一种坚持信念而缺少怀疑的结果。因此，尽管我们非常善于提出诸如新的指标，用于检验的新法则等新的想法，但是人类在处理糟糕的想法等问题时仍然无能为力。此外，我们人类的智慧还没有进化到对极度复杂或者随机的现象进行解释的程度，这揭示了我们在进行结构性思考时所暴露的局限性，因此，当我们需要把大量的变量转换成为独立的预测和判断时，对我们来说就是一项充满智慧的工作了。

虽然计算机本身并不具备发明创新能力，但是数据挖掘软件可以让它实现处理彼此相关性不强的指标，以及将复杂的模型进行合并的目的，而复杂模型是指能够把大量的变量转换成为具有预测性的事物。换句话说，计算机同样可以进行结构性思考。计算机的这种将复杂的高维模型进行合成的能力仅仅受到可得数据总量的局限。我们把对具有预测性的模型的约束称为“维数灾难”，它是由数学家理查德·贝尔曼（Richard Bellman）命名的。从本质上来讲，随着由多维空间组成的指标（每个指标由一个维度代表）的数量的增加，填入空间的观察结果就会以指数倍的速度降低密度，我们需要用达到一定水平的数据密度，通过数据挖掘软件把有效的形态从

随机背景中分辨出来。为了保持给定水平的数据密度，那么随着维度的增加，我们对观察资料数量的增加也将会以指数倍增长。如果在二维层面上能够符合数据密度的观察资料是 100 个，那么在三维层面上所需的观察资料就是 1 000 个，而四维层面上需要的观察资料就是 10 000 个。尽管这些在数据挖掘中受到实际的限制，但是现代数据挖掘软件在结构性推理和发现复杂形态方面的能力，要远强于最聪明的人类。

从本质上讲，计算机可以做人类不能做的事情，而人类也可以做计算机不能做的事情。人机交互真的是一种完美结合。

作用有限的主观预测者

尽管人机交互的协同作用具有一定的潜力，但是许多传统的技术分析实践者仍然坚持使用主观的方法从大量的指标当中获得预测或者信号。他们这样做，面对的是在过去的 50 年当中所积累的大量实证，而这些实证则论证了该方法的无用性。这些研究显示，反复出现的预测问题横跨各个领域，主观预测要比以简单的统计学法则（模型）为基础做出的预测差得多。重复的预测问题是指，相同类型的预测一定会在类似的信息集合的基础上被重复地做出。重复做出预测的实例包括技术分析、预测被释放的囚犯的暴力行为以及预测借款人是否会出现违约等。决策制定者可以获得的信息集合其实就是对大量的变量进行阅读。在技术分析领域，就是对一组技术指标进行解读；在信用评估领域，就是对财务比率进行分析。这么做的目的就是，确定指标和将要预测的结果之间是否存在稳定的函数关系，以及确定这些联系的本质是什么，然后再根据函数关系进行变量的合并。当这种关系非常复杂且包含更多的因素时，决策的制定者就要面对超出人类智慧极限的结构性思维。主观的专家们现在唯一的选择就是回归直觉判断。

在一项由保罗·米尔（Paul Meehl）³³于 1954 年进行的开创性研究中，他把主观预测（专家的直觉）与基于统计学法则（模型）做出的预测进行了对比。这项被称为“元分析”的研究是对过去的 20 项研究进行的再检验，其中，对依靠主观诊断的心理学家以及通过线性统计学模型得出诊断的精神病医生进行了比较。这些研究囊括了惯犯再次犯罪的可能性、预测电休

克疗法的结果等学术成就做出的预测结果。在每一个案例中，专家通过主观方式对大量的变量进行评估，然后做出一个判断。在所有的研究中，统计学模型提供了更加准确的预测。³⁴ 由索耶³⁵ 进行的后续研究对 45 个研究进行了元分析。在对全球的临床判断中，没有一项单独研究的表现能够超过统计学预测（索耶称其为“有机组合”）。³⁶ 索耶的研究值得我们注意，他认为，人类专家可以使用统计学模型所没有考虑到的信息，结果，模型的表现仍然非常优秀。例如，37 500 名在海军训练学校进行训练的水手的学习成绩，能够通过以测验成绩为基础建立的模型做出更加准确的预测，而以测验成绩以及进行个人面试做出的判断则达不到准确的预测。很明显，最后的判断对他们认为通过面试获得的微妙信息的想法过于重视，而对客观的可以进行衡量的事实视而不见。

嘉梅尔³⁷ 的一篇论文对专家的主观预测的准确性和两种客观的方法进行了比较：通过回归分析推导出的多变量线性模型和以专家做出判断的数学模型为基础的方法。准确性最低的是由专家做出的主观判断，而准确性最高的是由线性回归模型做出的预测。图 9-9 显示了由卢梭和休马克³⁸ 对嘉梅尔得出的结果的总结。用于比较预测法则的品质因数是预测的结果与实际的结果之间的相关系数。相关系数的范围可以设定为从 0 ~ 1.0。0 表示预测是毫无价值的，1.0 表示预测完美准确。预测问题横跨 9 个不同的领域：①研究生的学习成绩；②癌症患者的预期寿命；③股票价格的变化；④利用性格测试精神疾病；⑤心理学导论课程中的成绩和态度；⑥使用财务比率分析企业倒闭；⑦学生评教；⑧营销人员销售人寿保险的绩效；⑨使用罗尔沙赫氏试验对 IQ 进行评分。统计学模型与专家的平均相关性之比是 0.64 : 0.33。根据由相关系数的平方或者拟合度进行衡量的信息平均是拥有同样信息的专家的 3.76 倍。

大量额外的研究将专家的判断与统计学模型（法则）进行的比较证实了这些发现，即当人们试图把大量的变量进行合并之后做出预测或者判断的时候，人们的表现是非常差劲的。1968 年，戈德伯格³⁹ 证明，把使用人格测试评分作为输入量的线性预测模型，要比临床诊断更好地对精神病患者和神经症患者进行区分。模型的准确度是 70%，而专家的准确度则介

于 52% 与 67% 之间。1893 年的一项研究显示, 当变量的数量不断增加的时候, 专家们的表现是多么的糟糕。这项研究的任务是在 19 个输入量的基础上, 预测新确诊的男性精神病人的暴力倾向。用专家对暴力行为的预测和实际表现出的暴力行为之间的相关系数来衡量专家的平均预测准确度, 而这个值只有可怜的 0.12。这些专家之中只有一位的的成绩最好, 达到了 0.36。使用了 19 个相同的输入量集合的线性统计学模型所做出的预测是 0.82。在这种情况下, 模型的预测力约是那些掌握大量信息的专家做出的预测的 50 倍。

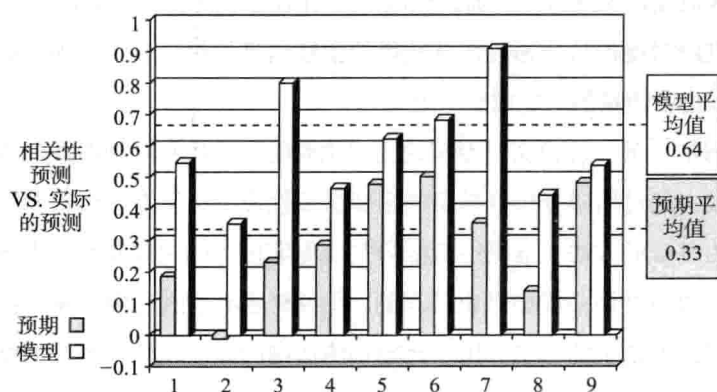


图 9-9 主观专家预测的绩效 VS. 9 个领域中客观的线性模型

米尔继续扩大着他对专家和统计学模型之间的比较的研究, 并且于 1986 年做出了如下的总结: “社会科学中是不存在争议的, 从如此庞大的多元化研究体系所得出的结果与通过社会科学而得出的结果是相同的。”你用 90 个观察资料 (现在要超过 150 个)⁴⁰ 就可以对足球比赛结果进行预测。⁴¹

实证继续在积累, 专家们也逐渐留意这些。关于对专家预测的准确性与线性统计学预测模型做出的预测进行比较的最全面的研究是由格罗夫和米尔于 1996 年完成的,⁴² 它包括 136 项研究, 其中有 96% 的法则等于或者超过专家的预测。在 2000 年的一项类似的研究中, Swets、Monahan 和 Dawes 对医学和心理学中的三种预测方法进行了比较: ①专家的主观判断; ②统计学法则 (模型); ③前两种方法的合并, 其中, 统计学法则的输出量

经过专家判断的修正。他们的发现与之前研究的结果是一样的，这就导致了继续按照主观的判断进行操作的专家并不仅仅是在预测的感觉上犯了错误，而且还会在开设处方药的时候犯错误。⁴³

综上所述，黑斯蒂（Hastie）和道斯（Dawes）⁴⁴做出了下列总结。

（1）专家和新手依靠规模相对较小的输入量（3～5个）做出主观判断。对牲畜的评估和天气预报是被排除在外的，这是大量的输入量被使用的原因。我们可以通过获得直接和准确的反馈这一事实对使用大量信息的能力进行解释。这也可以让两个领域的专家从中学到些东西。而在临床诊断、入学申请以及对金融市场进行预报等领域，反馈通常会出现滞后的现象，或者根本不能获得。黑斯蒂和道斯的结论提出了主观的技术分析师不能从他们的错误中吸取教训的原因，即反馈是不可获得的，这是因为缺少对客观形态的定义和评估标准。

（2）对于大范围的判断工作来说，专家的判断可以通过额外的（线性的）输入模型进行解释。

（3）虽然他们都以为自己是这样做的，但只有少数判断者使用非线性/结构性思维将输入量进行合并。然而，当他们进行结构性推理的时候，结果往往非常差。⁴⁵

（4）专家们并不知道自己是如何形成自己的判断的。这种情况在那些经验丰富的专家身上表现的更加突出。结果是，当把不同的输入量应用到不同的场合的时候，他们的判断就会出现不一致。

（5）在许多领域，当面对同样的信息集合时，每个人的判断都会不同于他人。由于相互间的判断不能达成高度的一致，所以他们就会说：有些人一定是搞错了。

（6）当把不相关的信息加入到输入量当中的时候，尽管最终的预测准确性没有得到提高，但是我们仍然充满自信。很明显，他们无法对相关的变量和不相关的变量进行区分。

（7）只有极少数判断者确实表现得非常突出。

黑斯蒂和道斯的总结得到了泰特洛克（Tetlock）的支持。⁴⁶泰特洛克对来自不同领域的专家的预测进行了评估，并且把它们与推断最近趋势的

相对简单的统计学模型进行了比较。泰特洛克的工作看起来似乎是对政治和经济领域里的专家做出的判断进行的最为严谨的长期研究，而他的研究又向我们揭示了这样一个事实，即制定决策的某些主观的认知类型是非常优秀的。他把这种认知的类型称作欺骗行为。然而，泰特洛克的研究显示，以人为基础做出的预测，不论其类型如何，在以正式模型为基础做出的预测面前，都显得是那么的苍白无力。使用二维评级框架评估预测的准确性、识别力和标准。以广义自回归滞后这种方法为基础建立的正式模型的预测准确性，要远远超过人类专家的表现。⁴⁷

为什么专家的主观预测表现如此差劲

尽管在那些对专家的判断进行研究的人当中，对专家的判断要比客观的预测差这个问题上达成了共识，但在关于是什么原因导致这种情况的发生这个问题上却存在着分歧。有人提出了以下几种解释：受到情感因素的困扰，在如何为系数进行加权时缺少统一性，对没有预测力的事物赋予过分的权重并使其具有生动具体的特点，赋予具有预测力的抽象的统计数据不充足的权重，没有正确使用结构性法则，没有被虚幻的相关性所迷惑，没有注意到有效的相关性以及受到他人言行的影响。本书第2章对上述内容都有涉及，关于错误观念的产生和坚持请参见第7章，至于信息瀑布和羊群效应，我就不在这里赘述了。我现在考虑的是主观的预测者受到情感因素的困扰问题，这个问题在之前的章节中还从未提到。

有些研究者认为主观判断在不确定的情况下才会受到情绪的影响。由洛文斯坦（Loewenstein）等人⁴⁸提出的一个模型指出，主观的预测受到决策者对预测结果的预期情绪反应的影响。虽然专家对于可能出现错误的焦虑在预测中并没有发挥有效的作用，但实际上这种焦虑确实发挥了作用。诺夫辛格（Nofsinger）⁴⁹用一张图表呈现了洛文斯坦的模型（见图9-10）。请注意，人们对可能出现的结果做出回应的情绪将会对预测者产生影响。这一点对认知性信息的评估和决策者⁵⁰当前的情绪状态都会产生影响。

斯洛维奇（Slovic）、菲纽肯（Finucane）、麦克格雷格⁵¹利用“影响”这个词解释了主观决策者受到的情绪困扰。影响是指因为受到刺激而产生

的好的或者坏的感觉。这种感觉的出现往往都是无意识的，并且在决策者的意识控制范围之外快速做出反射行为。斯洛维奇指出，感觉和想象对判断具有很大的影响，而这种影响可以通过对实证进行纯粹的分析和理性的思考这两种方式得到解决。例如，对不断下跌的市场抱有负面情绪的市场分析师仍然会比经过证明的实证更加持有看涨的观点。然而，并不是所有基于影响做出的判断都是不好的。某些情况确实体现着它们所引发的情绪，并且要求人们根据这些思想做出快速的决定。情感决定也许是人类进化过程中占据主导地位的判断模式，换句话说，做出情感决定的速度是具有生存价值的。由于对恐高症（指股民既想时时抄底，又不愿顺涨势而为，股价涨了，就拼命寻找低价股的行为）和大型捕食者（指对其他企业进行恶意收购的公司）的畏惧，我们就需要比那些做事之前权衡利弊的人对此进行更加快速的选择。然而，与其他提供一般适应性反应，又偶尔使我们误入歧途的启发式一样，过分地依靠“影响”也会让我们受骗上当。⁵² 当一种情况的相关特征并不能够由其激发的情绪很好地表现出来时，这种情况就发生了。金融市场就属于上述的这种情况。然而，主观预测仍然对基于影响的启发式缺少无意识的和潜意识的判断。

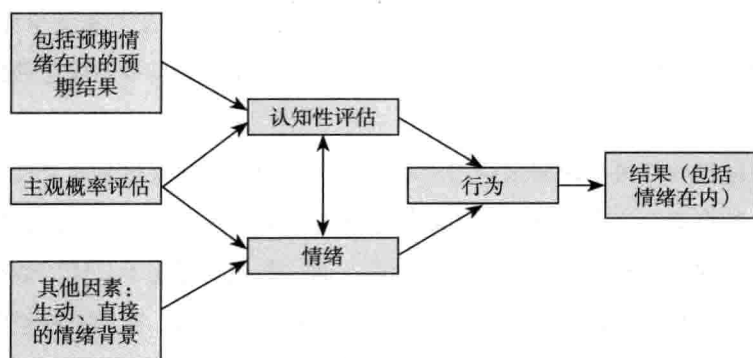


图 9-10 受到情感影响的洛文斯坦判断模型

由约瑟夫·福加斯（Joseph Forgas）提出的模型对主观预测的无理性做出了进一步的阐述。⁵³ 该模型认为，环境的复杂性和不确定性越大，影响决策的情绪就越多。⁵⁴ 当试图用一种理性的方式对多个指标进行评估时，主观的技术分析师将要面临的问题就是处理这些难以逾越的障碍。我个人

主张, 主观的预测将不再是分析师所扮演的角色。

专家在人机交互中发挥特有的和重要的作用

技术分析的实践者在设想的人机交互中所发挥的特有作用有两个:
①提供适合数据挖掘软件演示的信息含量丰富的备选输入量; ②对将要由数据挖掘软件解决的问题进行说明(如目标变量)。换句话说, 技术分析专家的特有作用就是提供计算机缺少的东西, 也就是说, 提供关于技术分析的专业知识和有创造力的人才。

具有讽刺意味的是, 在这个高科技领域, 分析师起到的作用成为了焦点。这一切都是改变了经济学的数据挖掘所导致的后果。在过去的 10 年里, 计算成本和强有力的数据挖掘软件的价格一落千丈。而这些资源在 20 年前则被资金充足的研究机构独占, 今天, 几乎所有的市场参与者都可以获得这份资源。这种趋势使得数据挖掘越来越平民化了。大多数投资者现在开始有权使用这些曾经为最大的机构投资者带来巨大竞争优势的工具了。现在, 这种巨大的优势已经被大大地削弱了。结果, 其竞争优势——附加值的主要来源也已经回归到分析师这一边了。如今, 那些在提出指标和对目标变量进行说明方面具有良好洞察力的专家更具备竞争优势。那些接受实证技术分析哲学思想的技术分析的实践者是具备发挥这种作用的能力的。

指标设计与目标说明是适用于本书的大课题, 现在有很多涉及技术分析指标的书籍,⁵⁵ 这里我就不再重复这些内容了。还有很多书籍讨论的是更加普遍的设计输入量问题(预测变量), 而这些问题适用于具体的数据挖掘算法。⁵⁶ 不同的算法需要具备的条件也不相同。例如, 虽然把神经网络中的输入量控制在合理的范围是非常重要的, 但是对于决策树算法来说就不必要了。马斯特斯(Masters)、⁵⁷ 派尔(Pyle)、⁵⁸ 魏斯(Weiss)及因度克亚(Indurkha)⁵⁹ 对这些问题进行了很好的论述。

指标设计是指在这些指标被提交给数据挖掘工具之前, 对应用到原始市场数据的指标进行具体的改造。数据挖掘器通常把这个过程称为“预处理”。要想完成为数据挖掘问题提供一组信息含量丰富且恰当的备选指标

这项任务，我们需要运用分析师的专业知识和创新天赋。

任何在数据挖掘战壕中弄脏双手的人都应当了解下面的真理：要想获得数据挖掘工作的成功，我们需要的是一套优秀的备选指标。“优秀”并不意味着所有的输入量都必须信息含量丰富，然而，有些被提出的变量则必须包含对目标变量来说非常有用的信息。多利安·派尔是数据挖掘领域的权威专家，他曾经说过：“在我们发现了需要解决的正确问题之后，数据准备通常是解决问题的关键所在。”我们可以很容易地发现成功与失败、有用的洞察力和令人费解的混沌、有价值的预测和无用的猜测之间的差异。⁶⁰另外一位数据挖掘领域的专家蒂莫斯·马斯特斯博士则认为：“预处理是通往成功的关键……如果变量通过预处理这种方法清楚地显示了重要信息，那么即使是最简单的预测模型也能获得上佳的绩效。”⁶¹在魏斯和因度克亚的著作 *Predictive Data Mining* 中，他们对信息进行了放大处理。通过描述更多的具有预测性的功能（指标），可以获得最大的绩效。只有人类才能够对一组特征的集合进行解释……也只有人类才知道如何把原有的功能转换为更好的功能。计算机非常善于删除无用的功能，但是在面对诸如整理新功能和把原始数据转化成更具预测力的形式等更高的要求时，往往又显示出滞后性，多种功能的组合对于结果的质量来说是更具决定性的因素，而对用于产生这些结果的具体预测方法来说，就不是决定性的因素了。在大多数情况下，功能的组合取决于对知识的应用。⁶²

理解数据挖掘方法及其在指标设计中的重要性的技术分析的实践者，将会在 21 世纪人机交互的技术分析中发挥这一重要作用。现在，计算机已经被作为智能放大器，并且在理念体系中进行了科学定位，实证技术分析的实践者在加速推动市场知识前沿的过程中将会具有独一无二的优势。对技术分析来说，鉴于人类的智慧不会发生太大的变化，而计算机智能正在以数倍的速度发展，⁶³ 已经没有任何其他更合理的方法了。

附录

Appendix

对消除趋势与以头寸偏差为基础的基准相等的观点进行论证

为了弥补市场偏差与多头－空头偏差相互作用产生的损失，我们需要用相同的多头－空头偏差来计算随机系统的预期收益率，并且从备选系统收益率中减去带有偏差的收益率。

根据定义，从历史原始收益率中提取的单一原始收益率的预期值就是这些收益率的平均值。

$$E_{\text{原始}} = \frac{1}{n} \sum_{i=1}^n R_i \quad (\text{A-1})$$

P_i 表示在时机 i 这一点上交易系统的头寸。此处的头寸可以是多头、空头或者中性。头寸处于随机状态下的交易系统的预期收益率就应当为

$$E_{\text{随机}} = \sum_{P_i=\text{多头}} E_{\text{原始}} - \sum_{P_i=\text{空头}} E_{\text{原始}} \quad (\text{A-2})$$

当备选系统持有多头头寸时，备选法则的全部收益率就是在那些时点上的原始收益率的总和，当备选系统持有空头头寸时，其全部原始收益率的总和就是负数

$$\text{总收益率} = \sum_{P_i=\text{多头}} R_i - \sum_{P_i=\text{空头}} R_i \quad (\text{A-3})$$

备选系统的总收益率（式（A-3））是通过减去具有相似偏差的随机系统的收益率得到的经过修正（式（A-2））的多头/空头头寸。那么经过修正的收益率就等于用对多头进行求和的值减去对空头进行求和的值

$$\text{修正的} = \sum_{P_i=\text{多头}} (R_i - E_{\text{原始}}) - \sum_{P_i=\text{空头}} (R_i - E_{\text{原始}}) \quad (\text{A-4})$$

请注意，经过修正的式（A-4）的收益率与没有经过修正的式（A-3）的收益率是相等的，二者间的不同之处在于要把原始收益率的平均值从每一个独立的原始收益率中减去。换句话说，为了消除存在偏差的市场中的多头 / 空头不平衡所导致的影响，我们要做的就是将原始收益率进行集中。

阿伦森教授用简洁的语言论证了在解释技术分析方法论的过程中存在的理论缺陷,即有效市场假说的假设与结论的缺陷,并且用适当的技术对技术分析的方法进行了改进与检验,事实证明这是正确的。读者将会从此书中受益匪浅。

——**杰克·施瓦格 (Jack Schwager)**
《金融奇才》(Market Wizards) 与“施瓦格期货丛书”的作者

阿伦森教授对数据挖掘的阐述是每一位分析师的必读之物,他对统计推断的全面论述是非常成功的。本书所选取的常识性案例不仅为检验与检验的确认过程节省了时间、金钱,而且还大大减少了在检验过程中遇到的困难。

——**佩里·考夫曼 (Perry Kaufman)**
《交易系统与方法》(New Trading Systems and Methods) 作者

本书打破了很多存在于技术分析领域的神话。读者在购买有关技术分析体系的书籍之前,应当先仔细阅读本书。

——**桑德尔·施特劳斯 (Sandor Straus)**
Merfin LLC公司高管

在这部充满争议而又引人注目的著作中,你可以不认同戴维·阿伦森的每一个观点。然而,对于每一个试图用科学的严谨态度进行投资技术分析的交易者来说,应当从辩证的角度来阅读本书。

——**纳尔逊·弗莱堡 (Nelson Freeburg)**
Formula Research编辑

Evidence-Based Technical Analysis

技术分析领域出现了一种新的趋势,即利用科学的程序,并依靠统计检验来衡量某一结论质量的新的、严谨的技术分析类型。正如定量分析长期使用的技术分析工具那样,市场分析也开始使用这一工具了。《实证技术分析》是一座里程碑,它将引导我们向着利用更好、更可靠的技术分析类型的目标前进。

——**约翰·布林格 (John Bollinger)**
CFA, CMT, Bollingerbands.com

戴维·阿伦森对当下流行的技术分析方法的批评可谓恰到好处,这能够让我们感到他的每一次批评都为我们厘清一次思维和视觉上的谬误。

——**弗雷德·杰姆 (Fred Gehr)**
《量化交易与资金管理》(Quantitative Trading and Money Management) 的作者

WILEY

www.wiley.com



收藏性: ★★★★★
实战性: ★★★★★
专业性: ★★★★★

阅读提示

入门

进阶

专业

投稿热线: (010) 88379007

客服热线: (010) 68995261 88361066

购书热线: (010) 68326294 88379649 68995259

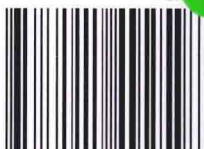
华章网站: www.hzbook.com

网上购书: www.china-pub.com

数字阅读: www.hzmedia.com.cn

上架指导: 投资/证券/股票

ISBN 978-7-111-49259-7



9 787111 492597 >

定价: 75.00元